

*Annual Review of Economics*

# Machine Learning Methods That Economists Should Know About

Susan Athey<sup>1,2,3</sup> and Guido W. Imbens<sup>1,2,3,4</sup>

<sup>1</sup>Graduate School of Business, Stanford University, Stanford, California 94305, USA;  
email: [athey@stanford.edu](mailto:athey@stanford.edu), [imbens@stanford.edu](mailto:imbens@stanford.edu)

<sup>2</sup>Stanford Institute for Economic Policy Research, Stanford University, Stanford,  
California 94305, USA

<sup>3</sup>National Bureau of Economic Research, Cambridge, Massachusetts 02138, USA

<sup>4</sup>Department of Economics, Stanford University, Stanford, California 94305, USA

**ANNUAL  
REVIEWS CONNECT**

[www.annualreviews.org](http://www.annualreviews.org)

- Download figures
- Navigate cited references
- Keyword search
- Explore related articles
- Share via email or social media

Annu. Rev. Econ. 2019. 11:685–725

First published as a Review in Advance on  
June 10, 2019

The *Annual Review of Economics* is online at  
[economics.annualreviews.org](http://economics.annualreviews.org)

<https://doi.org/10.1146/annurev-economics-080217-053433>

Copyright © 2019 by Annual Reviews.  
All rights reserved

JEL code: C30

## Keywords

machine learning, causal inference, econometrics

## Abstract

We discuss the relevance of the recent machine learning (ML) literature for economics and econometrics. First we discuss the differences in goals, methods, and settings between the ML literature and the traditional econometrics and statistics literatures. Then we discuss some specific methods from the ML literature that we view as important for empirical researchers in economics. These include supervised learning methods for regression and classification, unsupervised learning methods, and matrix completion methods. Finally, we highlight newly developed methods at the intersection of ML and econometrics that typically perform better than either off-the-shelf ML or more traditional econometric methods when applied to particular classes of problems, including causal inference for average treatment effects, optimal policy estimation, and estimation of the counterfactual effect of price changes in consumer choice models.

## 1. INTRODUCTION

In the abstract of his provocative 2001 paper in *Statistical Science*, the Berkeley statistician Leo Breiman (2001b, p. 199) writes about the difference between model-based and algorithmic approaches to statistics:

There are two cultures in the use of statistical modeling to reach conclusions from data. One assumes that the data are generated by a given stochastic data model. The other uses algorithmic models and treats the data mechanism as unknown.

Breiman (2001b, p. 199) goes on to claim that,

The statistical community has been committed to the almost exclusive use of data models. This commitment has led to irrelevant theory, questionable conclusions, and has kept statisticians from working on a large range of interesting current problems. Algorithmic modeling, both in theory and practice, has developed rapidly in fields outside statistics. It can be used both on large complex data sets and as a more accurate and informative alternative to data modeling on smaller data sets. If our goal as a field is to use data to solve problems, then we need to move away from exclusive dependence on data models and adopt a more diverse set of tools.

Breiman's (2001b) characterization no longer applies to the field of statistics. The statistics community has by and large accepted the machine learning (ML) revolution that Breiman refers to as the algorithm modeling culture, and many textbooks discuss ML methods alongside more traditional statistical methods (e.g., Hastie et al. 2009, Efron & Hastie 2016). Although the adoption of these methods in economics has been slower, they are now beginning to be widely used in empirical work and are the topic of a rapidly increasing methodological literature. In this review, we want to make the case that economists and econometricians also, as Breiman writes about the statistics community, "need to move away from exclusive dependence on data models and adopt a more diverse set of tools." We discuss some of the specific tools that empirical researchers would benefit from, and that we feel should be part of the standard graduate curriculum in econometrics if, as Breiman writes and we agree with, "our goal as a field is to use data to solve problems;" if, in other words, we view econometrics as, in essence, decision making under uncertainty (e.g., Chamberlain 2000); and if we wish to enable students to communicate effectively with researchers in other fields where these methods are routinely being adopted. Although relevant more generally, the methods developed in the ML literature have been particularly successful in big data settings, where we observe information on a large number of units, many pieces of information on each unit, or both, and often outside the simple setting with a single cross-section of units. For such settings, ML tools are becoming the standard across disciplines, so the economist's toolkit needs to adapt accordingly while preserving the traditional strengths of applied econometrics.

Why has the acceptance of ML methods been so much slower in economics compared to the broader statistics community? A large part of it may be the culture as Breiman refers to it. Economics journals emphasize the use of methods with formal properties of a type that many of the ML methods do not naturally deliver. This includes large sample properties of estimators and tests, including consistency, normality, and efficiency. In contrast, the focus in the ML literature is often on working properties of algorithms in specific settings, with the formal results being of a different type, e.g., guarantees of error rates. There are typically fewer theoretical results of the type traditionally reported in econometrics papers, although recently there have been some major advances in this area (Wager & Athey 2017, Farrell et al. 2018). There are no formal results that show that, for supervised learning problems, deep learning or neural net methods are uniformly superior to regression trees or random forests, and it appears unlikely that general results for such comparisons will soon be available, if ever.

Although the ability to construct valid large-sample confidence intervals is important in many cases, one should not out-of-hand dismiss methods that cannot deliver them (or, possibly, that cannot yet deliver them) if these methods have other advantages. The demonstrated ability to outperform alternative methods on specific data sets in terms of out-of-sample predictive power is valuable in practice, even though such performance is rarely explicitly acknowledged as a goal or assessed in econometrics. As Mullainathan & Spiess (2017) highlight, some substantive problems are naturally cast as prediction problems, and assessing their goodness of fit on a test set may be sufficient for the purposes of the analysis in such cases. In other cases, the output of a prediction problem is an input to the primary analysis of interest, and statistical analysis of the prediction component beyond convergence rates is not needed. However, there are also many settings where it is important to provide valid confidence intervals for a parameter of interest, such as an average treatment effect. The degree of uncertainty captured by standard errors or confidence intervals may be a component in decisions about whether to implement the treatment. We argue that, in the future, as ML tools are more widely adopted, researchers should articulate clearly the goals of their analysis and why certain properties of algorithms and estimators may or may not be important.

A major theme of this review is that, even though there are cases where using simple off-the-shelf algorithms from the ML literature can be effective (for examples, see Mullainathan & Spiess 2017), there are also many cases where this is not the case. The ML techniques often require careful tuning and adaptation to effectively address the specific problems that economists are interested in. Perhaps the most important type of adaptation is to exploit the structure of the problems, e.g., the causal nature of many estimands; the endogeneity of variables; the configuration of data such as panel data; the nature of discrete choice among a set of substitutable products; or the presence of credible restrictions motivated by economic theory, such as monotonicity of demand in prices or other shape restrictions (Matzkin 1994, 2007). Statistics and econometrics have traditionally put much emphasis on these structures and developed insights to exploit them, whereas ML has often put little emphasis on them. Exploitation of these insights, both substantive and statistical, which, in a different form, is also seen in the careful tuning of ML techniques for specific problems such as image recognition, can greatly improve their performance. Another type of adaptation involves changing the optimization criteria of ML algorithms to prioritize considerations from causal inference, such as controlling for confounders or discovering treatment effect heterogeneity. Finally, techniques such as sample splitting [using different data to select models than to estimate parameters (e.g., Athey & Imbens 2016, Wager & Athey 2017)] and orthogonalization (e.g., Chernozhukov et al. 2016a) can be used to improve the performance of ML estimators, in some cases leading to desirable properties such as asymptotic normality of ML estimators (e.g., Athey et al. 2016b, Farrell et al. 2018).

In this review, we discuss a list of tools that we feel should be part of the empirical economist's toolkit and should be covered in the core econometrics graduate courses. Of course, this is a subjective list, and given the speed with which this literature is developing, the list will rapidly evolve. Moreover, we do not give a comprehensive discussion of these topics; rather, we aim to provide an introduction to these methods that conveys the main ideas and insights, with references to more comprehensive treatments. First on our list is nonparametric regression, or in the terminology of the ML literature, supervised learning for regression problems. Second, we discuss supervised learning for classification problems or, closely related but not quite the same, nonparametric regression for discrete response models. This is the area where ML methods have had perhaps their biggest successes. Third, we discuss unsupervised learning, or clustering analysis and density estimation. Fourth, we analyze estimates of heterogeneous treatment effects and optimal policies mapping from individuals' observed characteristics to treatments. Fifth,

we discuss ML approaches to experimental design, where bandit approaches are starting to revolutionize effective experimentation, especially in online settings. Sixth, we discuss the matrix completion problem, including its application to causal panel data models and problems of consumer choice among a discrete set of products. Finally, we discuss the analysis of text data.

We note that there are a few other recent reviews of ML methods aimed at economists, often with more empirical examples and references to applications than we discuss in this review. Varian (2014) provides an early high-level discussion of a selection of important ML methods. Mullainathan & Spiess (2017) focus on the benefits of supervised learning methods for regression and discuss the prevalence of problems in economics where prediction methods are appropriate. Athey (2017) and Athey et al. (2017c) provide a broader perspective with more emphasis on recent developments in adapting ML methods for causal questions and general implications for economics. Gentzkow et al. (2017) provide an excellent recent discussion of methods for text analyses with a focus on economics applications. In the computer science and statistics literatures, there are also several excellent textbooks, with different levels of accessibility to researchers with a social science background, including the work of Efron & Hastie (2016); Hastie et al. (2009), who provide a more comprehensive text from a statistics perspective; Burkov (2019), who provides a very accessible introduction; Alpaydin (2009); and Knox (2018); all of these works take more of a computer science perspective.

## 2. ECONOMETRICS AND MACHINE LEARNING: GOALS, METHODS, AND SETTINGS

In this section, we introduce some of the general themes of this review. What are the differences in the goals and concerns of traditional econometrics and the ML literature, and how do these goals and concerns affect the choices among specific methods?

### 2.1. Goals

The traditional approach in econometrics, as exemplified in leading texts such as those of Greene (2000), Angrist & Pischke (2008), and Wooldridge (2010), is to specify a target, an estimand, that is a functional of a joint distribution of the data. The target is often a parameter of a statistical model that describes the distribution of a set of variables (typically conditional on some other variables) in terms of a set of parameters, which can be a finite or infinite set. Given a random sample from the population of interest, the parameter of interest and the nuisance parameters are estimated by finding the parameter values that best fit the full sample, using an objective function such as the sum of squared errors or the likelihood function. The focus is on the quality of the estimators of the target, traditionally measured through large sample efficiency. There is often also interest in constructing confidence intervals. Researchers typically report point estimates and standard errors.

In contrast, in the ML literature, the focus is typically on developing algorithms [a widely cited paper by Wu et al. (2008) has the title “Top 10 Algorithms in Data Mining”]. The goal for the algorithms is typically to make predictions about some variables given others or to classify units on the basis of limited information, for example, to classify handwritten digits on the basis of pixel values.

In a very simple example, suppose that we model the conditional distribution of some outcome  $Y_i$  given a vector-valued regressor or feature  $\mathbf{X}_i$ . Suppose that we are confident that

$$Y_i|\mathbf{X}_i \sim \mathcal{N}(\alpha + \boldsymbol{\beta}^\top \mathbf{X}_i, \sigma^2).$$

We could estimate  $\theta = (\alpha, \beta)$  by least squares, that is, as

$$(\hat{\alpha}_{ls}, \hat{\beta}_{ls}) = \arg \min_{\alpha, \beta} \sum_{i=1}^N (Y_i - \alpha - \beta^\top \mathbf{X}_i)^2.$$

Most introductory econometrics texts would focus on the least squares estimator without much discussion. If the model is correct, then the least squares estimator has well-known attractive properties: It is unbiased, it is the best linear unbiased estimator, it is the maximum likelihood estimator, and thus it has large sample efficiency properties.

In ML settings, the goal may be to make a prediction for the outcome for new units on the basis of their regressor values. Suppose that we are interested in predicting the value of  $Y_{N+1}$  for a new unit  $N+1$ , on the basis of the regressor values for this new unit,  $\mathbf{X}_{N+1}$ . Suppose that we restrict ourselves to linear predictors, so that the prediction is

$$\hat{Y}_{N+1} = \hat{\alpha} + \hat{\beta}^\top \mathbf{X}_{N+1}$$

for some estimator  $(\hat{\alpha}, \hat{\beta})$ . The loss associated with this decision may be the squared error

$$(Y_{N+1} - \hat{Y}_{N+1})^2.$$

The question now is how to come up with estimators  $(\hat{\alpha}, \hat{\beta})$  that have good properties associated with this loss function. This need not be the least squares estimator. In fact, when the dimension of the features exceeds two, we know from decision theory that we can do better in terms of expected squared error than the least squares estimator. The latter is not admissible; that is, there are other estimators that dominate the least squares estimator.

## 2.2. Terminology

One source of confusion is the use of new terminology in ML for concepts that have well-established labels in the older literatures. In the context of a regression model, the sample used to estimate the parameters is often referred to as the training sample. Instead of the model being estimated, it is being trained. Regressors, covariates, or predictors are referred to as features. Regression parameters are sometimes referred to as weights. Prediction problems are divided into supervised learning problems, where we observe both the predictors (features)  $\mathbf{X}_i$  and the outcome  $Y_i$ , and unsupervised learning problems, where we only observe the  $\mathbf{X}_i$  and try to group them into clusters or otherwise estimate their joint distribution. Unordered discrete response problems are generally referred to as classification problems.

## 2.3. Validation and Cross-Validation

In most discussions of linear regression in econometric textbooks, there is little emphasis on model validation. The form of the regression model, be it parametric or nonparametric, and the set of regressors are assumed to be given from the outside, e.g., economic theory. Given this specification, the task of the researcher is to estimate the unknown parameters of this model. Much emphasis is placed on doing this estimation step efficiently, typically operationalized through definitions of large sample efficiency. If there is discussion of model selection, it is often in the form of testing null hypotheses concerning the validity of a particular model, with the implication that there is a true model that should be selected and used for subsequent tasks.

Consider the regression example in the previous section. Let us assume that we are interested in predicting the outcome for a new unit, randomly drawn from the same population as our sample

was drawn from. As an alternative to estimating the linear model with an intercept and a scalar  $\mathbf{X}_i$ , we could estimate the model with only an intercept. Certainly, if  $\beta = 0$ , then that model would lead to better predictions. By the same argument, if the true value of  $\beta$  were close but not exactly equal to zero, then we would still do better leaving  $\mathbf{X}_i$  out of the regression. Out-of-sample cross-validation can help guide such decisions. There are two components of the problem that are important for this ability. First, the goal is predictive power, rather than estimation of a particular structural or causal parameter. Second, the method uses out-of-sample comparisons, rather than in-sample goodness-of-fit measures. This ensures that we obtain unbiased comparisons of the fit.

## 2.4. Overfitting, Regularization, and Tuning Parameters

The ML literature is much more concerned with overfitting than the standard statistics or econometrics literatures. Researchers attempt to select flexible models that fit well, but not so well that out-of-sample prediction is compromised. There is much less emphasis on formal results that particular methods are superior in large samples (asymptotically); instead, methods are compared on specific data sets to see what works well. A key concept is that of regularization. As Vapnik (2013, p. 9) writes, “Regularization theory was one of the first signs of the existence of intelligent inference.”

Consider a setting with a large set of models that differ in their complexity, measured, for example, as the number of unknown parameters in the model or, more subtly, through the the Vapnik–Chervonenkis (VC) dimension that measures the capacity or complexity of a space of models. Instead of directly optimizing an objective function, say, minimizing the sum of squared residuals in a least squares regression setting or maximizing the logarithm of the likelihood function, a term is added to the objective function to penalize the complexity of the model. There are antecedents of this practice in the traditional econometrics and statistics literatures. One is that, in likelihood settings, researchers sometimes add a term to the logarithm of the likelihood function equal to minus the logarithm of the sample size times the number of free parameters divided by two, leading to the Bayesian information criterion, or simply the number of free parameters, the Akaike information criterion. In Bayesian analyses of regression models, the use of a prior distribution on the regression parameters, centered at zero, independent across parameters with a constant prior variance, is another way of regularizing estimation that has a long tradition. The modern approaches to regularization are different in that they are more data driven, with the amount of regularization determined explicitly by the out-of-sample predictive performance rather than by, for example, a subjectively chosen prior distribution.

Consider a linear regression model with  $K$  regressors,

$$Y_i | \mathbf{X}_i \sim \mathcal{N}(\beta^\top \mathbf{X}_i, \sigma^2).$$

Suppose that we also have a prior distribution for the the slope coefficients  $\beta_k$ , with the prior for  $\beta_k$ ,  $\mathcal{N}(0, \tau^2)$ , and independent of  $\beta_{k'}$  for any  $k \neq k'$ . (This may be more plausible if we first normalize the features and outcome to have mean zero and unit variance. We assume that this has been done.) Given the value for the variance of the prior distribution,  $\tau^2$ , the posterior mean for  $\beta$  is the solution to

$$\arg \min_{\beta} \sum_{i=1}^N (Y_i - \beta^\top \mathbf{X}_i)^2 + \frac{\sigma^2}{\tau^2} \|\beta\|_2^2,$$

where  $\|\beta\|_2 = (\sum_{k=1}^K \beta_k^2)^{1/2}$ . One version of an ML approach to this problem is to estimate  $\beta$  by minimizing

$$\arg \min_{\beta} \sum_{i=1}^N (Y_i - \beta^T \mathbf{X}_i)^2 + \lambda \|\beta\|_2^2.$$

The only difference is in the way the penalty parameter  $\lambda$  is chosen. In a formal Bayesian approach, this reflects the (subjective) prior distribution on the parameters, and it would be chosen a priori. In an ML approach,  $\lambda$  would be chosen through out-of-sample cross-validation to optimize the out-of-sample predictive performance. This is closer to an empirical Bayes approach, where the data are used to estimate the prior distribution (e.g., Morris 1983).

## 2.5. Sparsity

In many settings in the ML literature, the number of features is substantial, both in absolute terms and relative to the number of units in the sample. However, there is often a sense that many of the features are of minor importance, if not completely irrelevant. The problem is that we may not know *ex ante* which of the features matter and which can be dropped from the analysis without substantially hurting the predictive power.

Hastie et al. (2009, 2015) discuss what they call the sparsity principle:

Assume that the underlying true signal is sparse and we use an  $\ell_1$  penalty to try to recover it. If our assumption is correct, we can do a good job in recovering the true signal.... But if we are wrong—the underlying truth is not sparse in the chosen bases—then the  $\ell_1$  penalty will not work well. However, in that instance, no method can do well, relative to the Bayes error. (Hastie et al. 2015, page 24)

Exact sparsity is in fact stronger than is necessary; in many cases it is sufficient to have approximate sparsity, where most of the explanatory variables have very limited explanatory power, even if not zero, and only a few of the features are of substantial importance (see, for example, Belloni et al. 2014).

Traditionally, in the empirical literature in social sciences, researchers limited the number of explanatory variables by hand, rather than choosing them in a data-dependent manner. Allowing the data to play a bigger role in the variable selection process appears to be a clear improvement, even if the assumption that the underlying process is at least approximately sparse is still a very strong one, and even if inference in the presence of data-dependent model selection can be challenging.

## 2.6. Computational Issues and Scalability

Compared to the traditional statistics and econometrics literatures, the ML literature is much more concerned with computational issues and the ability to implement estimation methods with large data sets. Solutions that may have attractive theoretical properties in terms of statistical efficiency but that do not scale well to large data sets are often discarded in favor of methods that can be implemented easily in very large data sets. This can be seen in the discussion of the relative merits of least absolute shrinkage and selection operator (LASSO) versus subset selection in linear regression settings. In a setting with a large number of features that might be included in the analysis, subset selection methods focus on selecting a subset of the regressors and then estimating the parameters of the regression function by least squares. However, LASSO has computational

advantages. It can be implemented by adding a penalty term that is proportional to the sum of the absolute values of the parameters. A major attraction of LASSO is that there are effective methods for calculating the LASSO estimates with the number of regressors in the millions. Best subset selection regression, in contrast, is an NP-hard problem. Until recently, it was thought that this was only feasible in settings with the number of regressors in the 30s, although current research (Bertsimas et al. 2016) suggests that it may be feasible with the number of regressors in the 1,000s. This has reopened a new, still unresolved debate on the relative merits of LASSO versus best subset selection (see Hastie et al. 2017) in settings where both are feasible. There are some indications that, in settings with a low signal-to-noise ratio, as is common in many social science applications, LASSO may have better performance, although there remain many open questions. In many social science applications, the scale of the problems is such that best subset selection is also feasible, and the computational issues may be less important than these substantive aspects of the problems.

A key computational optimization tool used in many ML methods is stochastic gradient descent (SGD) (Bottou 1998, 2012; Friedman 2002). It is used in a wide variety of settings, including in optimizing neural networks and estimating models with many latent variables (e.g., Ruiz et al. 2017). The idea is very simple. Suppose that the goal is to estimate a parameter  $\theta$ , and that the estimation approach entails finding the value  $\hat{\theta}$  that minimizes an empirical loss function, where  $Q_i(\theta)$  is the loss for observation  $i$ , and the overall loss is the sum  $\sum_i Q_i(\hat{\theta}_k)$ , with derivative  $\sum_i \nabla Q_i(\hat{\theta})$ . Classic gradient descent methods involve an iterative approach, where  $\hat{\theta}_k$  is updated from  $\hat{\theta}_{k-1}$  as follows:

$$\theta_k = \theta_{k-1} - \eta_k \frac{1}{N} \sum_i \nabla Q_i(\hat{\theta}),$$

where  $\eta_k$  is the learning rate, often chosen optimally through line search. More sophisticated optimization methods multiply the first derivative by the inverse of the matrix of second derivatives or estimates thereof.

The challenge with this approach is that it can be computationally expensive. The computational cost is in evaluating the full derivative  $\sum_i \nabla Q_i$  and even more in optimizing the learning rate  $\eta_k$ . The idea behind SGD is that it is better to take many small steps that are noisy but, on average, in the right direction than it is to spend equivalent computational cost in very accurately figuring out in what direction to take a single small step. More specifically, SGD uses the fact that the average of  $\nabla Q_i$  for a random subset of the sample is an unbiased (but noisy) estimate of the gradient. For example, dividing the data randomly into 10 subsets or batches, with  $B_i \in \{1, 10\}$  denoting the subset unit  $i$  belongs to, one could do 10 steps of the type

$$\theta_k = \theta_{k-1} - \eta_k \frac{1}{N/10} \sum_{i:B_i=k} \nabla Q_i(\hat{\theta}_k),$$

with a deterministic learning rate  $\eta_k$ . After the 10 iterations, one could reshuffle the data set and then repeat. If the learning rate  $\eta_k$  decreases at an appropriate rate, then under relatively mild assumptions, SGD converges almost surely to a global minimum when the objective function is convex or pseudoconvex and otherwise converges almost surely to a local minimum. Bottou (2012) provides an overview and practical tips for implementation.

The idea can be pushed even further in the case where  $\nabla Q_i(\theta)$  is itself an expectation. We can consider evaluating  $\nabla Q_i$  using Monte Carlo integration. However, rather than taking many Monte Carlo draws to get an accurate approximation to the integral, we can instead take a small number



of draws or even a single draw. This type of approximation is used in economic applications by Ruiz et al. (2017) and Hartford et al. (2016).

## 2.7. Ensemble Methods and Model Averaging

Another key feature of the ML literature is the use of model averaging and ensemble methods (e.g., Dietterich 2000). In many cases, a single model or algorithm does not perform as well as a combination of possibly quite different models, averaged using weights (sometimes called votes) obtained by optimizing out-of-sample performance. A striking example is the Netflix Prize competition (Bennett & Lanning 2007), where all of the top contenders use combinations of models and often averages of many models (Bell & Koren 2007). There are two related ideas in the traditional econometrics literature. Obviously, Bayesian analysis implicitly averages over the posterior distribution of the parameters. Mixture models are also used to combine different parameter values in a single prediction. However, in both cases, this model averaging involves averaging over similar models, typically with the same specification, that are only different in terms of parameter values. In the modern literature, and in the top entries in the Netflix Prize competition, the models that are averaged over can be quite different, and the weights are obtained by optimizing out-of-sample predictive power, rather than in-sample fit.

For example, one may have three predictive models, one based on a random forest, leading to predictions  $\hat{Y}_i^{\text{RF}}$ ; one based on a neural net, with predictions  $\hat{Y}_i^{\text{NN}}$ ; and one based on a linear model estimated by LASSO, leading to  $\hat{Y}_i^{\text{LASSO}}$ . Then, using a test sample, one can choose weights  $p^{\text{rf}}$ ,  $p^{\text{nn}}$ , and  $p^{\text{lasso}}$  by minimizing the sum of squared residuals in the test sample:

$$(\hat{p}^{\text{rf}}, \hat{p}^{\text{nn}}, \hat{p}^{\text{lasso}}) = \arg \min_{p^{\text{rf}}, p^{\text{nn}}, p^{\text{lasso}}} \sum_{i=1}^{N^{\text{test}}} \left( Y_i - p^{\text{rf}} \hat{Y}_i^{\text{RF}} - p^{\text{nn}} \hat{Y}_i^{\text{NN}} - p^{\text{lasso}} \hat{Y}_i^{\text{LASSO}} \right)^2,$$

$$\text{subject to } p^{\text{rf}} + p^{\text{nn}} + p^{\text{lasso}} = 1 \quad \text{and} \quad p^{\text{rf}}, p^{\text{nn}}, p^{\text{lasso}} \geq 0.$$

One may also estimate weights based on regression of the outcomes in the test sample on the predictors from the different models without imposing that the weights sum to one and are non-negative because random forests, neural nets, and LASSO have distinct strengths and weaknesses in terms of how well they deal with the presence of irrelevant features, nonlinearities, and interactions. As a result, averaging over these models may lead to out-of-sample predictions that are strictly better than predictions based on a single model.

In a panel data context (Athey et al. 2019), one can use ensemble methods combining various forms of synthetic control and matrix completion methods and find that the combinations outperform the individual methods.

## 2.8. Inference

The ML literature has focused heavily on out-of-sample performance as the criterion of interest. This has come at the expense of one of the concerns that the statistics and econometrics literatures have traditionally focused on, namely, the ability to do inference, e.g., construct confidence intervals that are valid, at least in large samples. Efron & Hastie (2016, p. 209) write:

Prediction, perhaps because of its model-free nature, is an area where algorithmic developments have run far ahead of their inferential justification.

Although there has recently been substantial progress in the development of methods for inference for low-dimensional functionals in specific settings [e.g., the work of Wager & Athey (2017) in the context of random forests and of Farrell et al. (2018) in the context of neural networks], it remains the case that, for many methods, it is currently impossible to construct confidence intervals that are valid, even if only asymptotically. One question is whether this ability to construct confidence intervals is as important as the traditional emphasis on it in the econometric literature suggests. For many decision problems, it may be that prediction is of primary importance, and inference is at best of secondary importance. Even in cases where it is possible to do inference, it is important to keep in mind that the requirements that ensure this ability often come at the expense of predictive performance. One can see this tradeoff in traditional kernel regression, where the bandwidth that optimizes expected squared error balances the tradeoff between the square of the bias and the variance, so that the optimal estimators have an asymptotic bias that invalidates the use of standard confidence intervals. This can be fixed by using a bandwidth that is smaller than the optimal one, so that the asymptotic bias vanishes, but it does so explicitly at the expense of increasing the variance.

### 3. SUPERVISED LEARNING FOR REGRESSION PROBLEMS

One of the canonical problems in both the ML and econometric literatures is that of estimating the conditional mean of a scalar outcome given a set of covariates or features. Let  $Y_i$  denote the outcome for unit  $i$ , and let  $\mathbf{X}_i$  denote the  $K$ -component vector of covariates or features. The conditional expectation is

$$g(x) = \mathbb{E}[Y_i | \mathbf{X}_i = x].$$

Compared to the traditional econometric textbooks (e.g., Greene 2000, Angrist & Pischke 2008, Wooldridge 2010), there are some conceptual differences in the ML literature (for discussion, see Mullainathan & Spiess 2017). In the settings considered in the ML literature, there are often many covariates, sometimes more than there are units in the sample. There is no presumption in the ML literature that the conditional distribution of the outcomes given the covariates follows a particular parametric model. The derivatives of the conditional expectation for each of the covariates, which in the linear regression model correspond to the parameters, are not of intrinsic interest. Instead, the focus is on out-of-sample predictions and their accuracy. Furthermore, there is less of a sense that the conditional expectation is monotone in each of the covariates compared to many economic applications. There is often concern that the conditional expectation may be an extremely nonmonotone function with some higher-order interactions of substantial importance.

The econometric literature on estimating the conditional expectation is also huge. Parametric methods for estimating  $g(\cdot)$  often use least squares. Since the work of Bierens (1987), kernel regression methods have become a popular alternative when more flexibility is required, and series or sieve methods have subsequently gained interest (for a survey, see Chen 2007). These methods have well-established large sample properties, allowing for the construction of confidence intervals. Simple nonnegative kernel methods are viewed as performing very poorly in settings with high-dimensional covariates, with the difference  $\hat{g}(x) - g(x)$  of order  $O_p(N^{-1/K})$ . This rate can be improved by using higher-order kernels and assuming the existence of many derivatives of  $g(\cdot)$ , but practical experience with high-dimensional covariates has not been satisfactory for these methods, and applications of kernel methods in econometrics are generally limited to low-dimensional settings.

The differences in performance between some of the traditional methods such as kernel regression and the modern methods such as random forests are particularly pronounced in sparse settings with a large number of more or less irrelevant covariates. Random forests are effective at picking up on the sparsity and ignoring the irrelevant features, even if there are many of them, while the traditional implementations of kernel methods essentially waste degrees of freedom on accounting for these covariates. Although it may be possible to adapt kernel methods for the presence of irrelevant covariates by allowing for covariate-specific bandwidths, in practice there has been little effort in this direction. A second issue is that the modern methods are particularly good at detecting severe nonlinearities and high-order interactions. The presence of such high-order interactions in some of the success stories of these methods should not blind us to the fact that, with many economic data, we expect high-order interactions to be of limited importance. If we try to predict earnings for individuals, then we expect the regression function to be monotone in many of the important predictors such as education and prior earnings variables, even for homogeneous subgroups. This means that models based on linearizations may do well in such cases relative to other methods, compared to settings where monotonicity is fundamentally less plausible, as, for example, in an image recognition problem. This is also a reason for the superior performance of locally linear random forests (Friedberg et al. 2018) relative to standard random forests.

We discuss four specific sets of methods, although there are many more, including variations on the basic methods. First, we discuss methods where the class of models considered is linear in the covariates, and the question is solely about regularization. Second, we discuss methods based on partitioning the covariate space using regression trees and random forests. Third, we discuss neural nets, which were the focus of a small econometrics literature in the 1990s (Hornik et al. 1989, White 1992) but more recently have become a very prominent part of the literature on ML in various subtle reincarnations. Fourth, we discuss boosting as a general principle.

### 3.1. Regularized Linear Regression: LASSO, Ridge, and Elastic Nets

Suppose that we consider approximations to the conditional expectation that have a linear form

$$g(x) = \beta^\top x = \sum_{k=1}^K \beta_k x_k,$$

after the covariates and the outcome are demeaned, and the covariates are normalized to have unit variance. The traditional method for estimating the regression function in this case is least squares, with

$$\hat{\beta}^{\text{ls}} = \arg \min_{\beta} \sum_{i=1}^N (Y_i - \beta^\top \mathbf{X}_i)^2.$$

However, if the number of covariates  $K$  is large relative to the number of observations  $N$ , then the least squares estimator  $\hat{\beta}_k^{\text{ls}}$  does not even have particularly good repeated sampling properties as an estimator for  $\beta_k$ , let alone good predictive properties. In fact, with  $K \geq 3$ , the least squares estimator is not even admissible and is dominated by estimators that shrink toward zero. With  $K$  very large, possibly even exceeding the sample size  $N$ , the least squares estimator has particularly poor properties, even if the conditional mean of the outcome given the covariates is in fact linear.

Even with  $K$  modest in magnitude, the predictive properties of the least squares estimator may be inferior to those of estimators that use some amount of regularization. One common form of

regularization is to add a penalty term that shrinks the  $\beta_k$  toward zero and minimize

$$\arg \min_{\beta} \sum_{i=1}^N (Y_i - \beta^\top \mathbf{X}_i)^2 + \lambda (\|\beta\|_q)^{1/q},$$

where  $\|\beta\|_q = \sum_{k=1}^K |\beta_k|^q$ . For  $q = 1$ , this corresponds to LASSO (Tibshirani 1996). For  $q = 2$ , this corresponds to ridge regression (Hoerl & Kennard 1970). As  $q \rightarrow 0$ , the solution penalizes the number of nonzero covariates, leading to best subset regression (Miller 2002, Bertsimas et al. 2016). In addition, there are many hybrid methods and modifications, including elastic nets, which combine penalty terms from LASSO and ridge (Zou & Hastie 2005); the relaxed LASSO, which combines least squares estimates from the subset selected by LASSO and the LASSO estimates themselves (Meinshausen 2007); least angle regression (Efron et al. 2004); the Dantzig selector (Candès & Tao 2007); and the non-negative garrotte (Breiman 1993).

There are a couple of important conceptual differences among these three special cases, subset selection, LASSO, and ridge regression (for a recent discussion, see Hastie et al. 2017). First, both best subset and LASSO lead to solutions with a number of the regression coefficients exactly equal to zero, a sparse solution. For the ridge estimator, in contrast, all of the estimated regression coefficients will generally differ from zero. It is not always important to have a sparse solution, and the variable selection that is implicit in these solutions is often overinterpreted. Second, best subset regression is computationally hard (NP-hard) and, as a result, not feasible in settings with  $N$  and  $K$  large, although progress has recently been made in this regard (Bertsimas et al. 2016). LASSO and ridge regression have a Bayesian interpretation. Ridge regression gives the posterior mean and mode under a normal model for the conditional distribution of  $Y_i$  given  $\mathbf{X}_i$ , and normal prior distributions for the parameters. LASSO gives the posterior mode given Laplace prior distributions. However, in contrast to formal Bayesian approaches, the coefficient  $\lambda$  on the penalty term is, in the modern literature, chosen through out-of-sample cross-validation, rather than subjectively through the choice of prior distribution.

### 3.2. Regression Trees and Forests

Regression trees (Breiman et al. 1984) and their extension, random forests (Breiman 2001a), have become very popular and effective methods for flexibly estimating regression functions in settings where out-of-sample predictive power is important. They are considered to have great out-of-the-box performance without requiring subtle tuning. Given a sample  $(\mathbf{X}_{i1}, \dots, \mathbf{X}_{iK}, Y_i)$ , for  $i = 1, \dots, N$ , the idea is to split the sample into subsamples and estimate the regression function within the subsamples simply as the average outcome. The splits are sequential and based on a single covariate  $\mathbf{X}_{ik}$  at a time exceeding a threshold  $c$ . Starting with the full training sample, consider a split based on feature or covariate  $k$  and threshold  $c$ . The sum of in-sample squared errors before the split is

$$Q = \sum_{i=1}^N (Y_i - \bar{Y})^2,$$

where

$$\bar{Y} = \frac{1}{N} \sum_{i=1}^N Y_i.$$

After a split based on covariate  $k$  and threshold  $c$ , the sum of in-sample squared errors is

$$Q(k, c) = \sum_{i: X_{ik} \leq c} (Y_i - \bar{Y}_{k,c,l})^2 + \sum_{i: X_{ik} > c} (Y_i - \bar{Y}_{k,c,r})^2,$$

where (with  $l$  and  $r$  denoting left and right)

$$\bar{Y}_{k,c,l} = \sum_{i: X_{ik} \leq c} Y_i / \sum_{i: X_{ik} \leq c} 1$$

and

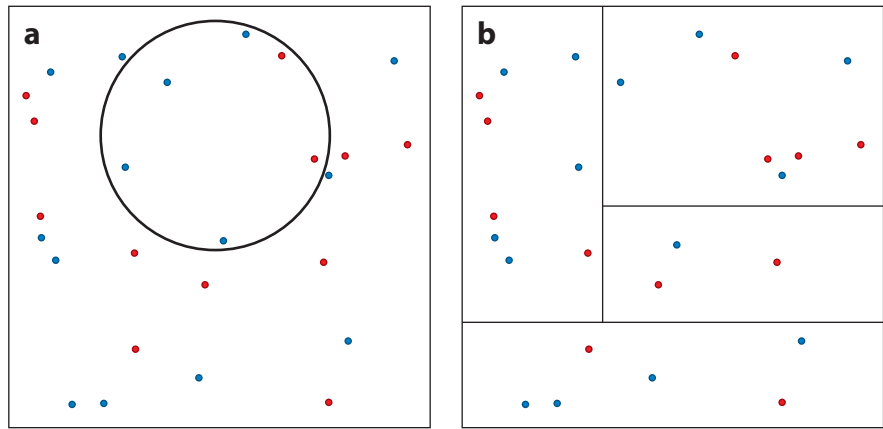
$$\bar{Y}_{k,c,r} = \sum_{i: X_{ik} > c} Y_i / \sum_{i: X_{ik} > c} 1$$

are the average outcomes in the two subsamples. We split the sample using the covariate  $k$  and threshold  $c$  that minimize the average squared error  $Q(k, c)$  over all covariates  $k = 1, \dots, K$  and all thresholds  $c \in (-\infty, \infty)$ . We then repeat this, optimizing also over the subsamples or leaves. At each split, the average squared error is further reduced (or stays the same). We therefore need some regularization to avoid the overfitting that would result from splitting the sample too many times. One approach is to add a penalty term to the sum of squared residuals that is linear in the number of subsamples (the leaves). The coefficient on this penalty term is then chosen through cross-validation. In practice, a very deep tree is estimated, and then pruned to a more shallow tree using cross-validation to select the optimal tree depth. The sequence of first growing and then pruning the tree avoids splits that may be missed because their benefits rely on subtle interactions.

An advantage of a single tree is that it is easy to explain and interpret results. Once the tree structure is defined, the prediction in each leaf is a sample average, and the standard error of that sample average is easy to compute. However, it is not, in general, true that the sample average of the mean within a leaf is an unbiased estimate of what the mean would be within that same leaf in a new test set. Since the leaves were selected using the data, the leaf sample means in the training data will tend to be more extreme (in the sense of being different from the overall sample mean) than in an independent test set. Athey & Imbens (2016) suggest sample splitting as a way to avoid this issue. If a confidence interval for the prediction is desired, then the analyst can simply split the data in half. One half of the data are used to construct a regression tree. Then, the partition implied by this tree is taken to the other half of the data, where the sample mean within a given leaf is an unbiased estimate of the true mean value for the leaf.

Although trees are easy to interpret, it is important not to go too far in interpreting the structure of the tree, including the selection of variables used for the splits. Standard intuitions from econometrics about omitted variable bias can be useful in this case. Particular covariates that have strong associations with the outcome may not show up in splits because the tree splits on covariates highly correlated with those covariates.

One way to interpret a tree is that it is an alternative to kernel regression. Within each tree, the prediction for a leaf is simply the sample average outcome within the leaf. Thus, we can think of the leaf as defining the set of nearest neighbors for a given target observation in a leaf, and the estimator from a single regression tree is a matching estimator with nonstandard ways of selecting the nearest neighbor to a target point. In particular, the neighborhoods will prioritize some covariates over others in determining which observations qualify as nearby. **Figure 1** illustrates the difference between kernel regression and a tree-based matching algorithm for the case of two



**Figure 1**

(a) Euclidean neighborhood for  $k$ -nearest neighbor (KNN) matching. (b) Tree-based neighborhood.

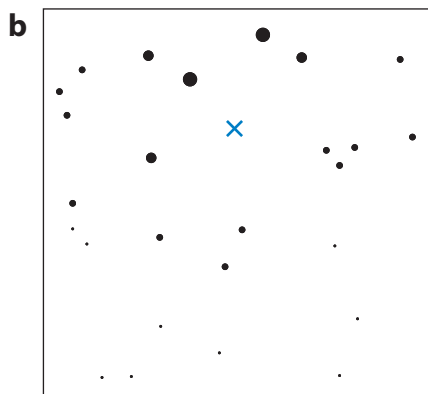
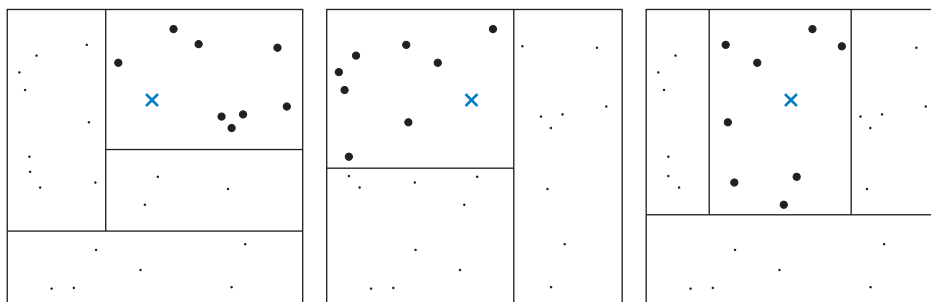
covariates. Kernel regression will create a neighborhood around a target observation based on the Euclidean distance to each point, while tree-based neighborhoods will be rectangles. In addition, a target observation may not be in the center of a rectangle. Thus, a single tree is generally not the best way to predict outcomes for any given test point  $x$ . When a prediction tailored to a specific target observation is desired, generalizations of tree-based methods can be used.

For better estimates of  $\mu(x)$ , random forests (Breiman 2001a) build on the regression tree algorithm. A key issue that random forests address is that the estimated regression function given a tree is discontinuous with substantial jumps, more so than one might like. Random forests induce smoothness by averaging over a large number of trees. These trees differ from each other in two ways. First, each tree is based not on the original sample, but on a bootstrap sample [known as bagging (Breiman 1996)] or, alternatively, on a subsample of the data. Second, the splits at each stage are not optimized over all possible covariates, but rather over a random subset of the covariates, changing every split. These two modifications lead to sufficient variation in the trees that the average is relatively smooth (although still discontinuous) and, more importantly, has better predictive power than a single tree.

Random forests have become very popular methods. A key attraction is that they require relatively little tuning and have great performance out of the box compared to more complex methods such as deep learning neural networks. Random forests and regression trees are particularly effective in settings with a large number of features that are not related to the outcome, that is, settings with sparsity. The splits will generally ignore those covariates, and as a result, the performance will remain strong even in settings with a large number of features. Indeed, when comparing forests to kernel regression, a reliable way to improve the relative performance of random forests is to add irrelevant covariates that have no predictive power. These will rapidly degrade the performance of kernel regression but will not affect a random forest nearly as severely because it will largely ignore them (Wager & Athey 2017).

Although the statistical analysis of forests has proved elusive since Breiman's original work, Wager & Athey (2017) show that a particular variant of random forests can produce estimates  $\hat{\mu}(x)$  with an asymptotically normal distribution centered on the true value  $\mu(x)$ ; furthermore, they provide an estimate of the variance of the estimator so that centered confidence intervals can be constructed. The variant that they study uses subsampling rather than bagging; furthermore,

**a** Different



**Figure 2**

(a) Different trees in a random forest generating weights for test point  $X$ . (b) The kernel based on the share of trees in the same leaf as test point  $X$ .

each tree is built using two disjoint subsamples, one used to define the tree and the second used to estimate sample means for each leaf. This honest estimation is crucial for the asymptotic analysis.

Random forests can be connected to traditional econometric methods in several ways. Returning to the kernel regression comparison, since each tree is a form of matching estimator, the forest is an average of matching estimators. As **Figure 2** illustrates, by averaging over trees, the prediction for each point will be centered on the test point (except near boundaries of the covariate space). However, the forest prioritizes more important covariates for selecting matches in a data-driven way. Another way to interpret random forests (e.g., Athey et al. 2016b) is that they generate weighting functions analogous to kernel weighting functions. For example, a kernel regression makes a prediction at a point  $x$  by averaging nearby points but weighting closer points more heavily. A random forest, by averaging over many trees, will include nearby points more often than distant points. We can formally derive a weighting function for a given test point by counting the share of trees where a particular observation is in the same leaf as a test point. Then, random forest predictions can be written as

$$\hat{\mu}_{\text{rf}}(x) = \sum_{i=1}^n \alpha_i(x) Y_i, \quad \sum_{i=1}^n \alpha_i(x) = 1, \quad \alpha_i(x) \geq 0, \quad 1.$$

where the weights  $\alpha_i(x)$  encode the weight given by the forest to the  $i$ th training example when predicting at  $x$ . The difference between typical kernel weighting functions and forest-based weighting

functions is that the forest weights are adaptive; if a covariate has little effect, it will not be used in splitting leaves, and thus the weighting function will not be very sensitive to distance along that covariate.

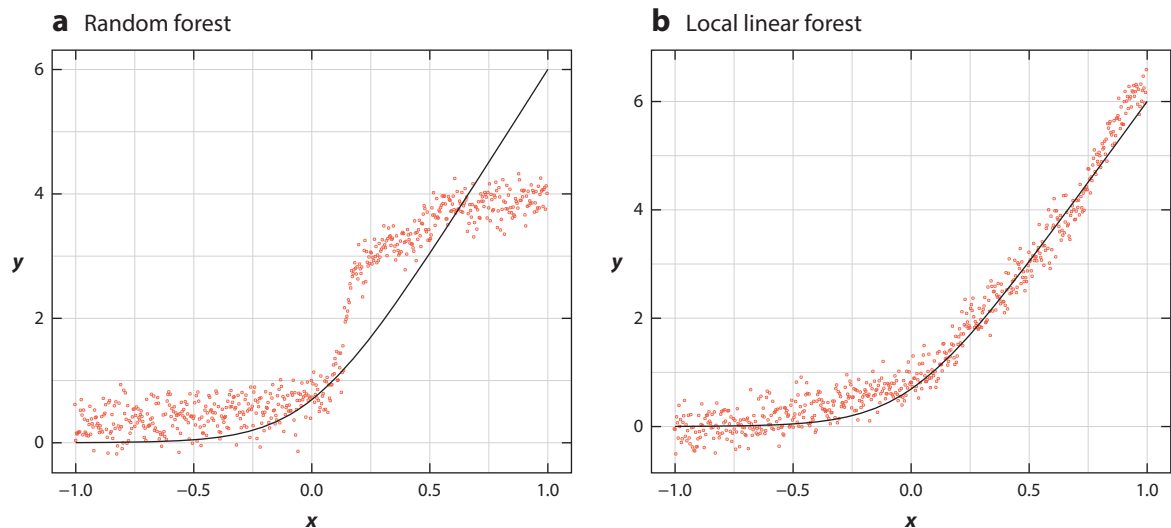
Recently random forests have been extended to settings where the interest is in causal effects, either average or unit-level (Wager & Athey 2017), as well as for estimating parameters in general economic models that can be estimated with maximum likelihood or generalized method of moments (GMM) (Athey et al. 2016b). In the latter case, the interpretation of the forest as creating a weighting function is operationalized; the new generalized random forest algorithm operates in two steps. First, a forest is constructed, and second, a GMM model is estimated for each test point, where points that are nearby in the sense of frequently occurring in the same leaf as the test point are weighted more heavily in estimation. With an appropriate version of honest estimation, these forests produce parameter estimates with an asymptotically normal distribution. Generalized random forests can be thought of as a generalization of local maximum likelihood, which was introduced by Tibshirani & Hastie (1987), but where kernel weighting functions are used to weight nearby observations more heavily than observations distant from a particular test point.

A weakness of forests is that they are not very efficient at capturing linear or quadratic effects or at exploiting smoothness of the underlying data-generating process. In addition, near the boundaries of the covariate space, they are likely to have bias because the leaves of the component trees of the random forest cannot be centered on points near the boundary. Traditional econometrics encounters this boundary bias problem in analyses of regression discontinuity designs where, for example, geographical boundaries of school districts or test score cutoffs determine eligibility for schools or programs (Imbens & Lemieux 2008). The solution proposed in the econometrics literature, for example, in the matching literature (Abadie & Imbens 2011), is to use local linear regression, which is a regression with nearby points weighted more heavily. Suppose that the conditional mean function is increasing as it approaches the boundary. Then, the local linear regression corrects for the fact that, at a test point near the boundary, most sample points lie in a region with lower conditional mean than the conditional mean at the boundary. Friedberg et al. (2018) extend the generalized random forest framework to local linear forests, which are constructed by running a regression weighted by the weighting function derived from a forest. In their simplest form, local linear forests just take the forest weights  $\alpha_i(x)$  and use them for local regression:

$$(\hat{\mu}(x), \hat{\theta}(x)) = \operatorname{argmin}_{\mu, \theta} \left\{ \sum_{i=1}^n \alpha_i(x) [Y_i - \mu(x) - (\mathbf{X}_i - x)\theta(x)]^2 + \lambda \|\theta(x)\|_2^2 \right\}. \quad 2.$$

Performance can be improved by modifying the tree construction to incorporate a regression correction; in essence, splits are optimized for predicting residuals from a local regression. This algorithm performs better than traditional forests in settings where a regression can capture broad patterns in the conditional mean function, such as monotonicity or a quadratic structure, and, again, asymptotic normality is established. **Figure 3** illustrates how local linear forests can improve on regular random forests: By fitting local linear regressions with a random forest–estimated kernel, the resulting predictions can match a simple polynomial function even in relatively small data sets. In contrast, a forest tends to have bias, particularly near boundaries, and in small data sets will have more of a step function shape. Although **Figure 3** shows the impact in a single dimension, an advantage of the forest over a kernel is that these corrections can occur in multiple dimensions while still allowing the traditional advantages of a forest of uncovering more complex interactions among covariates.





**Figure 3**

Predictions from random forests and local linear forests on 600 test points. Training and test data were simulated from  $Y_i = \log(1 + e^{\delta X_{i1}}) + \epsilon_i \sim N(0, 20)$ , with  $X$  having dimension  $d = 20$  (19 covariates are irrelevant) and errors  $\epsilon \sim N(0, 20)$ . Forests were trained on  $n = 600$  training points using the R package GRF and tuned via cross-validation. The true conditional mean signal  $\mu(x)$  is shown in black, and predictions are shown in red. Figure adapted with permission from Friedberg et al. (2018).

### 3.3. Deep Learning and Neural Nets

Using neural networks and related deep learning methods is another general and flexible approach to estimating regression functions. They have been found to be very successful in complex settings with extremely large numbers of features. However, in practice, these methods require a substantial amount of tuning to work well for a given application relative to methods such as random forests. Neural networks were studied in the econometric literature in the 1990s but did not catch on at the time (see Hornik et al. 1989, White 1992).

Let us consider a simple example. Given  $K$  covariates (features)  $\mathbf{X}_{ik}$ , we model  $K_1$  latent or unobserved variables  $Z_{ik}$  (hidden nodes) that are linear in the original covariates:

$$Z_{ik}^{(1)} = \sum_{j=1}^K \beta_{kj}^{(1)} \mathbf{X}_{ij} \text{ for } k = 1, \dots, K_1.$$

We then modify these linear combinations using a simple nonlinear transformation, e.g., a sigmoid function

$$g(z) = [1 + \exp(-z)]^{-1}$$

or a rectified linear function

$$g(z) = z \mathbf{1}_{z>0},$$

and then model the outcome as a linear function of this nonlinear transformation of these hidden nodes plus noise:

$$Y_i = \sum_{k=1}^{K_1} \beta_k^{(2)} g[Z_{ik}^{(1)}] + \varepsilon_i.$$

This is a neural network with a single hidden layer with  $K_1$  hidden nodes. The transformation  $g(\cdot)$  introduces nonlinearities in the model. Even with this single layer, with many nodes, one can approximate arbitrarily well a rich set of smooth functions.

It may be tempting to fit this into a standard framework and interpret this model simply as a complex, but fully parametric, specification for the potentially nonlinear conditional expectation of  $Y_i$  given  $\mathbf{X}_i$ :

$$\mathbb{E}[Y_i | \mathbf{X}_i = x] = \sum_{k'=1}^{K_1} \beta_{k'}^{(2)} g \left[ \sum_{k=1}^K \beta_{k'k}^{(1)} \mathbf{X}_{ik} \right].$$

Given this interpretation, we can estimate the unknown parameters using nonlinear least squares. We could then derive the properties of the least squares estimators, and functions thereof, under standard regularity conditions. However, this interpretation of a neural net as a standard nonlinear model would be missing the point, for four reasons. First, it is likely that the asymptotic distributions for the parameter estimates would be poor approximations to the actual sampling distributions. Second, the estimators for the parameters would be poorly behaved, with likely substantial collinearity without careful regularization. Third, and more important, these properties are not of intrinsic interest. We are interested in the properties of the predictions from these specifications, and these can be quite attractive even if the properties of the parameter estimates are not. Fourth, we can make these models much more flexible, and at the same time make the properties of the corresponding least squares estimators of the parameters substantially less tractable and attractive, by adding layers to the neural network. A second layer of hidden nodes would have representations that are linear in the same transformation  $g(\cdot)$  of linear combinations of the first layer of hidden nodes:

$$Z_{ik}^{(2)} = \sum_{j=1}^{K_1} \beta_{kj}^{(2)} g[Z_{ij}^{(1)}], \text{ for } k = 1, \dots, K_2,$$

with the outcome now a function of the second layer of hidden nodes,

$$Y_i = \sum_{k=1}^{K_2} \beta_k^{(3)} g[Z_{ik}^{(2)}] + \varepsilon_i.$$

The depth of the network substantially increases the flexibility in practice, even if, with a single layer and many nodes, we can already approximate a very rich set of functions. Asymptotic properties for multilayer networks have recently been established by Farrell et al. (2018). In applications, researchers have used models with many layers, e.g., ten or more, and millions of parameters:

We observe that shallow models [models with few layers] in this context overfit at around 20 millions parameters while deep ones can benefit from having over 60 million. This suggests that using a deep model expresses a useful preference over the space of functions the model can learn. (LeCun et al. 2015, p. 289)

In cases with multiple hidden layers and many hidden nodes, one needs to carefully regularize the parameter estimation, possibly through a penalty term that is proportional to the sum of the squared coefficients in the linear parts of the model. The architecture of the networks is also important. It is possible, as in the specification above, to have the hidden nodes at a particular layer be a linear function of all the hidden nodes of the previous layer, or to restrict them to a subset based on substantive considerations (e.g., proximity of covariates in some metric, such as location of pixels in a picture). Such convolutional networks have been very successful but require even more careful tuning (Krizhevsky et al. 2012).

Estimation of the parameters of the network is based on approximately minimizing the sum of the squared residuals, plus a penalty term that depends on the complexity of the model. This minimization problem is challenging, especially in settings with multiple hidden layers. The algorithms of choice use the back-propagation algorithm and variations thereon (Rumelhart et al. 1986) to calculate the exact derivatives with respect to the parameters of the unit-level terms in the objective function. These algorithms exploit in a clever way the hierarchical structure of the layers and the fact that each parameter enters only into a single layer. The algorithms then use stochastic gradient descent (Bottou 1998, 2012; Friedman 2002), described in Section 2.6, as a computationally efficient method for finding the approximate optimum.

### 3.4. Boosting

Boosting is a general-purpose technique to improve the performance of simple supervised learning methods (for a detailed discussion, see Schapire & Freund 2012). Let us say that we are interested in prediction of an outcome given a substantial number of features. Suppose that we have a very simple algorithm for prediction, a simple base learner. For example, we could have a regression tree with three leaves, that is, a regression tree based on two splits, where we estimate the regression function as the average outcome in the corresponding leaf. Such an algorithm on its own would not lead to a very attractive predictor in terms of predictive performance because it uses at most two of the many possible features. Boosting improves this base learner in the following way. Take for all units in the training sample the residual from the prediction based on the simple three-leaf tree model,  $Y_i - \hat{Y}_i^{(1)}$ . Now we apply the same base learner (in this case, the two-split regression tree) with the residuals as the outcome of interest (and with the same set of original features). Let  $\hat{Y}_i^{(2)}$  denote the prediction from combining the first and second steps. Given this new tree, we can calculate the new residual,  $Y_i - \hat{Y}_i^{(2)}$ . We can then repeat this step, using the new residual as the outcome and again constructing a two-split regression tree. We can do this many times and get a prediction based on reestimating the basic model many times on the updated residuals.

If we base our boosting algorithm on a regression tree with  $L$  splits, then it turns out that the resulting predictor can approximate any regression function that can be written as the sum of functions of  $L$  of the original features at a time. So, with  $L = 1$ , we can approximate any function that is additive in the features, and with  $L = 2$ , we can approximate any function that is additive in functions of the original features that allow for general second-order effects.

Boosting can also be applied using base learners other than regression trees. The key is to choose a base learner that is easy to apply many times without running into computational problems.

## 4. SUPERVISED LEARNING FOR CLASSIFICATION PROBLEMS

Classification problems are the focus of the other main branch of the supervised learning literature. The problem is, given a set of observations on a vector of features  $\mathbf{X}_i$  and a label  $Y_i$  (an

unordered discrete outcome), the goal is a function that assigns new units, on the basis of their features, to one of the labels. This is very closely related to discrete choice analysis in econometrics, where researchers specify statistical models that imply a probability that the outcome takes on a particular value, conditional on the covariates (features). Given such a probability, it is, of course, straightforward to predict a unique label, namely the one with the highest probability. However, there are differences between the two approaches. An important one is that, in the classification literature, the focus is often solely on the classification, the choice of a single label. One can classify given a probability for each label, but one does not need such a probability to do the classification. Many of the classification methods do not, in fact, first estimate a probability for each label, and so are not directly relevant in settings where such a probability is required. A practical difference is that the classification literature has often focused on settings where, ultimately, the covariates allow one to assign the label with almost complete certainty, as opposed to settings where even the best methods have high error rates.

The classic example is that of digit recognition. Based on a picture, coded as a set of, say, 16 or 256 black and white pixels, the challenge is to classify the image as corresponding to one of the ten digits from 0 to 9. In this case, ML methods have been spectacularly successful. Support vector machines (SVMs) (Cortes & Vapnik 1995) greatly outperformed other methods in the 1990s. More recently, deep convolutional neural networks (Krizhevsky et al. 2012) have improved error rates even further.

#### 4.1. Classification Trees and Forests

Trees and random forests are easily modified from a focus on estimation of regression functions to classification tasks (for a general discussion, see Breiman et al. 1984). Again, we start by splitting the sample into two leaves, based on a single covariate exceeding or not exceeding a threshold. We optimize the split over the choice of covariate and the threshold. The difference between the regression case and the classification case is in the objective function that measures the improvement from a particular split. In classification problems, this is called the impurity function. It measures, as a function of the shares of units in a given leaf with a particular label, how impure that particular leaf is. If there are only two labels, then we could simply assign the labels the numbers zero and one, interpret the problem as one of estimating the conditional mean, and use the average squared residual as the impurity function. That does not generalize naturally to the multilabel case. Instead, a more common impurity function, as a function of the  $M$  shares  $p_1, \dots, p_M$ , is the Gini impurity,

$$I(p_1, \dots, p_M) = - \sum_{m=1}^M p_m \ln(p_m).$$

This impurity function is minimized if the leaf is pure, meaning that all units in that leaf have the same label, and is maximized if the shares are all equal to  $1/M$ . The regularization typically works, again, through a penalty term on the number of leaves in the tree. The same extension from a single tree to a random forest that is discussed above for the regression case works for the classification case.

#### 4.2. Support Vector Machines and Kernels

SVMs (Vapnik 2013, Scholkopf & Smola 2001) make up another flexible set of methods for classification analyses. SVMs can also be extended to regression settings but are more naturally introduced in a classification context, and, for simplicity, we focus on the case with two possible labels.

Suppose that we have a set with  $N$  observations on a  $K$ -dimensional vector of features  $\mathbf{X}_i$  and a binary label  $Y_i \in \{-1, 1\}$  (we could use 0/1 labels, but using  $-1/1$  is more convenient). Given a  $K$ -vector of weights  $\omega$  (what we would typically call the parameters) and a constant  $b$  (often called the bias in the SVM literature), define the hyperplane  $x \in \mathbb{R}$  such that  $\omega^\top x + b = 0$ . We can think of this hyperplane defining a binary classifier  $\text{sgn}(\omega^\top \mathbf{X}_i + b)$ , with units  $i$  with  $\omega^\top x + b \geq 0$  classified as 1 and units with  $\omega^\top x + b < 0$  classified as  $-1$ . Now consider for each hyperplane [that is, for each pair  $(\omega, b)$ ] the number of classification errors in the sample. If we are very fortunate, then there would be some hyperplanes with no classification errors. In that case, there are typically many such hyperplanes, and we choose the one that maximizes the distance to the closest units. There will typically be a small set of units that have the same distance to the hyperplane (the same margin). These are called the support vectors.

We can write this as an optimization problem as

$$(\hat{\omega}, \hat{b}) = \arg \min_{\omega, b} \|\omega\|^2, \quad \text{subject to } Y_i(\omega^\top \mathbf{X}_i + b) \geq 1, \quad \text{for all } i = 1, \dots, N,$$

with classifier

$$\text{sgn}(\hat{\omega}^\top \mathbf{X}_i + \hat{b}).$$

Note that, if there is a hyperplane with no classification errors, then a standard logit model would not have a maximum likelihood estimator: The argmax of the likelihood function would diverge.

We can also write this problem in terms of the Lagrangian, with  $\alpha_i$  being the Lagrangian multiplier for the restriction  $Y_i(\omega^\top \mathbf{X}_i + b) \geq 1$ ,

$$\min_{\alpha, \omega, b} \left\{ \frac{1}{2} \|\omega\|^2 - \sum_{i=1}^N \alpha_i (Y_i(\omega^\top \mathbf{X}_i + b) - 1) \right\}, \quad \text{subject to } 0 \leq \alpha_i.$$

After concentrating out the weights  $\omega$ , this is equivalent to

$$\max_{\alpha} \left\{ \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j Y_i Y_j \mathbf{X}_i^\top \mathbf{X}_j \right\}, \quad \text{subject to } 0 \leq \alpha_i, \quad \sum_{i=1}^N \alpha_i Y_i = 0,$$

where  $\hat{b}$  solves  $\sum_i \hat{\alpha}_i [Y_i(\mathbf{X}_i^\top \omega + \hat{b}) - 1] = 0$ , with classifier

$$f(x) = \text{sgn} \left( \hat{b} + \sum_{i=1}^N Y_i \hat{\alpha}_i \mathbf{X}_i^\top x \right).$$

In practice, of course, we are typically in a situation where there exists no hyperplane without classification errors. In that case, there is no solution, as the  $\alpha_i$  diverge for some  $i$ . We can modify the classifier by adding the constraint that the  $\alpha_i \leq C$ . Scholkopf & Smola (2001) recommend setting  $C = 10N$ .

This is still a linear problem, differing from a logistic regression only in terms of the loss function. Units far away from the hyperplane do not affect the estimator as much in the SVM approach as they do in a logistic regression, leading to more robust estimates. However, the real power of the SVM approach is in the nonlinear case. We can think of this in terms of constructing a number of functions of the original covariates,  $\phi(\mathbf{X}_i)$ , and then finding the optimal hyperplane in the transformed feature space. However, because the features enter only through the inner

product  $\mathbf{X}_i^\top \mathbf{X}_j$ , it is possible to skip the step of specifying the transformations  $\phi(\cdot)$  and instead directly write the classifier in terms of a kernel  $K(x, z)$ , through

$$\max_{\alpha} \sum_{i=1}^N \left\{ \alpha_i - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j Y_i Y_j K(\mathbf{X}_i, \mathbf{X}_j) \right\}, \quad \text{subject to } 0 \leq \alpha_i \leq C, \quad \sum_{i=1}^N \alpha_i Y_i = 0,$$

where  $\hat{b}$  solves  $\sum_i \hat{\alpha}_i [Y_i(\mathbf{X}_i^\top \omega + \hat{b}) - 1] = 0$ , with classifier

$$f(x) = \text{sgn} \left[ \sum_{i=1}^N Y_i \alpha_i K(\mathbf{X}_i, x) + b \right].$$

Common choices for the kernel are  $k_b(x, z) = \exp(-(x - z)^\top(x - z)/b)$  and  $k_{\kappa, \Theta}(x, z) = \tanh(\kappa(x - z)^\top(x - z) + \Theta)$ . The parameters of the kernel, capturing the amount of smoothing, are typically chosen through cross-validation.

## 5. UNSUPERVISED LEARNING

A second major topic in the ML literature is unsupervised learning. In this case, we have several cases without labels. We can think of this as having several observations on covariates without an outcome. We may be interested in partitioning the sample into subsamples, or clusters, or in estimating the joint distribution of these variables.

### 5.1. *K*-Means Clustering

In this case, the goal is, given a set of observations on features  $\mathbf{X}_i$ , to partition the feature space into subspaces. These clusters may be used to create new features based on subspace membership. For example, we may wish to use the partitioning to estimate parsimonious models within each of the subspaces. We may also wish to use cluster membership as a way to organize the sample into types of units that may receive different exposures to treatments. This is an unusual problem, in the sense that there is no natural benchmark to assess whether a particular solution is a good one relative to some other one. A closely related approach that is more traditional in the econometrics and statistics literatures is the mixture model, where the distribution that generated the sample is modeled as a mixture of different distributions. The mixture components are similar in nature to the clusters.

A key method is the *K*-means algorithm (Hartigan & Wong 1979, Alpaydin 2009). Consider the case where we wish to partition the feature space into  $K$  subspaces or clusters. We wish to choose centroids  $b_1, \dots, b_K$  and then assign units to the cluster based on their proximity to the centroids. The basic algorithm works as follows. We start with a set of  $K$  centroids,  $b_1, \dots, b_K$ , elements of the feature space, and sufficiently spread out over this space. Given a set of centroids, assign each unit to the cluster that minimizes the distance between the unit and the centroid of the cluster:

$$C_i = \arg \min_{c \in \{1, \dots, K\}} \|\mathbf{X}_i - b_c\|^2.$$

Then update the centroids as the average of the  $\mathbf{X}_i$  in each of the clusters:

$$b_c = \sum_{i:C_i=c} \mathbf{X}_i / \sum_{i:C_i=c} 1.$$

Repeatedly iterate between the two steps. Choosing the number of clusters  $K$  is difficult because there is no direct cross-validation method to assess the performance of one value versus the other. This number is often chosen on substantive grounds, rather than in a data-driven way.

There are a large number of alternative unsupervised methods, including topic models, which we discuss below in Section 9. Unsupervised variants of neural nets are particularly popular for images and videos.

## 5.2. Generative Adversarial Networks

Let us consider the problem of estimation of a joint distribution given observations on  $\mathbf{X}_i$  for a random sample of units. A recent ML approach to this is generative adversarial networks (GANs) (Goodfellow et al. 2014, Arjovsky & Bottou 2017). The idea is to find an algorithm to generate data that look like the sample  $\mathbf{X}_1, \dots, \mathbf{X}_N$ . A key insight is that there is an effective way of assessing whether the algorithm is successful that is like a Turing test. If we have a successful algorithm, then we should not be able to tell whether data were generated by the algorithm or came from the original sample. Thus, we can assess the algorithm by training a classifier on data from the algorithm and a subsample from the original data. If the algorithm is successful, then the classifier will not be able to successfully classify the data as coming from the original data or the algorithm. The GAN then uses the relative success of the classification algorithm to improve the algorithm that generates the data, in effect pitting the classification algorithm against the generating algorithm.

This type of algorithm may also be an effective way of choosing simulation designs intended to mimic real-world data.

## 6. MACHINE LEARNING AND CAUSAL INFERENCE

An important difference between much of the econometrics literature and the ML literature is that the econometrics literature is often focused on questions beyond simple prediction. In many, arguably most, cases, researchers are interested in average treatment effects or other causal or structural parameters (for surveys, see Imbens & Wooldridge 2009, Abadie & Cattaneo 2018). Covariates that are of limited importance for prediction may still play an important role in estimating such structural parameters.

### 6.1. Average Treatment Effects

A canonical problem is that of estimating average treatment effects under unconfoundedness (Rosenbaum & Rubin 1983, Imbens & Rubin 2015). Given data on an outcome  $Y_i$ , a binary treatment  $W_i$ , and a vector of covariates or features  $\mathbf{X}_i$ , a common estimand, the average treatment effect (ATE), is defined as  $\tau = \mathbb{E}[Y_i(1) - Y_i(0)]$ , where  $Y_i(w)$  is the potential outcome that unit  $i$  would have experienced if their treatment assignment had been  $w$ . Under the unconfoundedness assumption, which ensures that potential outcomes are independent of the treatment assignment conditional on covariates

$$W_i \perp\!\!\!\perp [Y_i(0), Y_i(1)] \mid \mathbf{X}_i,$$

the ATE is identified. The ATE can be characterized in several different ways as a functional of the joint distribution of  $(W_i, \mathbf{X}_i, Y_i)$ . Three important ones are (a) as the covariate-adjusted difference between the two treatment groups, (b) as a weighted average of the outcomes, and (c) in terms of the influence or efficient score function

$$\begin{aligned}\tau &= \mathbb{E} [\mu(1, \mathbf{X}_i) - \mu(0, \mathbf{X}_i)] \\ &= \mathbb{E} \left[ \frac{Y_i W_i}{e(\mathbf{X}_i)} - \frac{Y_i (1 - W_i)}{1 - e(\mathbf{X}_i)} \right] \\ &= \mathbb{E} \left\{ \frac{[Y_i - \mu(1, \mathbf{X}_i)] W_i}{e(\mathbf{X}_i)} - \frac{[Y_i - \mu(0, \mathbf{X}_i)] (1 - W_i)}{1 - e(\mathbf{X}_i)} \right\} + \mathbb{E} [\mu(1, \mathbf{X}_i) - \mu(0, \mathbf{X}_i)],\end{aligned}\quad 3.$$

where

$$\mu(w, x) = \mathbb{E}[Y_i | W_i = w, \mathbf{X}_i = x]$$

and

$$e(x) = \mathbb{E}[W_i | \mathbf{X}_i = x].$$

One can estimate the ATE using the first representation by estimating the conditional outcome expectations  $\mu(\cdot)$ , using the second representation by estimating the propensity score  $e(\cdot)$ , or using the third representation by estimating both the conditional outcome expectation and the propensity score. Given a particular choice of representation, there is the question of the appropriate estimator for the particular conditional expectations that enter into that representation. For example, if we wish to use the first representation and want to consider linear models, it may seem natural to use LASSO or subset selection. However, as illustrated by Belloni et al. (2014), such a strategy could have very poor properties. The set of features that is optimal for inclusion when the objective is estimating  $\mu(\cdot)$  is not necessarily optimal for estimating  $\tau$ . The reason is that omitting from the regression covariates that are highly correlated with the treatment  $W_i$  can introduce substantial biases even if their correlation with the outcome is only modest. Thus, optimizing model selection solely for predicting outcomes is not the best approach. Belloni et al. (2014) propose using a covariate selection method that selects both covariates that are predictive of the outcome and covariates that are predictive of the treatment, and show that this substantially improves the properties of the corresponding estimator for  $\tau$ .

More recent methods focus on combining estimation of the conditional outcome expectations  $\mu(\cdot)$  with the propensity score  $e(\cdot)$  flexibly and in doubly robust methods (Robins & Rotnitzky 1995; Chernozhukov et al. 2016a,b) and methods that combine estimation of the conditional outcome expectations  $\mu(\cdot)$  with covariate balancing (Athey et al. 2016a). Covariate balancing is inspired by another common approach in ML, which frames data analysis as an optimization problem. In this case, instead of trying to estimate a primitive object, the propensity score  $e(\cdot)$ , the optimization procedure directly optimizes weights for the observations that lead to the same mean values of covariates in the treatment and control groups (Zubizarreta 2015). This approach allows for efficient estimation of ATEs even when the propensity score is too complex to estimate well. Because traditional propensity score weighting entails dividing by the estimated propensity score, instability in propensity score estimation can lead to high variability in estimates for the ATE. Furthermore, in an environment with many potential confounders, estimating the propensity score using regularization may lead to the omission of weak confounders that still contribute to bias. Directly optimizing for balancing weights can be more effective in environments with many weak confounders.



The case of estimating ATEs under unconfoundedness is an example of a more general theme from econometrics; typically, economists prioritize precise estimates of causal effects above predictive power (for further elaboration of this point, see Athey 2017, 2019). In instrumental variables models, it is common that goodness of fit falls by a substantial amount between an ordinary least squares regression and the second stage of a two-stage least squares model. However, the instrumental variables estimate of causal effects can be used to answer questions of economic interest, so the loss of predictive power is viewed as less important.

## 6.2. Orthogonalization and Cross-Fitting

A theme that has emerged across multiple distinct applications of ML to parameter estimation is that both practical performance and theoretical guarantees can be improved by using two simple techniques, both involving nuisance parameters that are estimated using ML. These can be illustrated through the lens of estimation of ATEs. Building from the third representation in Equation 3, the influence function of an efficient semiparametric estimator is

$$\psi(y, w, x) = \mu(1, x) - \mu(0, x) + \frac{w}{e(x)}[y - \mu(1, x)] + \frac{1 - w}{1 - e(x)}[y - \mu(0, x)],$$

with  $\Psi_i = \psi(Y_i, W_i, \mathbf{X}_i)$ . An estimate of the ATE can be constructed by first constructing estimates  $\hat{\mu}(w, x)$  and  $\hat{e}(x)$  and plugging those in to get an estimate  $\hat{\Psi}_i = \hat{\psi}(Y_i, W_i, \mathbf{X}_i)$  for each observation. Then, the sample average of  $\hat{\Psi}_i$  is an estimator for the ATE. This approach is analyzed by Bickel et al. (1998) and Van der Vaart (2000) for the general semiparametric case and by Chernozhukov et al. (2017) for the ATE case. A key result is that an estimator based on this approach is efficient if the estimators are sufficiently accurate in the following sense:

$$\mathbb{E}\{[\hat{\mu}(w, \mathbf{X}_i) - \mu(w, \mathbf{X}_i)]^2\}^{\frac{1}{2}} \mathbb{E}\{[\hat{e}(\mathbf{X}_i) - e(\mathbf{X}_i)]^2\}^{\frac{1}{2}} = o(N^{-1/2}).$$

For example, each nuisance component,  $\hat{\mu}(\cdot)$  and  $\hat{e}(\cdot)$ , could converge at a rate close to  $N^{-1/4}$ , an order of magnitude slower than the ATE estimate. This works because  $\Psi_i$  makes use of orthogonalization; by construction, errors in estimating the nuisance components are orthogonal to errors in  $\Psi_i$ . This idea is more general and has been exploited in a series of papers, with theoretical analysis discussed by Chernozhukov et al. (2018a,c) and other applications for estimating heterogeneous effects in models with unconfoundedness or those that make use of instrumental variables, as discussed by authors including Athey et al. (2016b).

A second idea, also exploited in the same series of papers, is that performance can be improved using techniques such as sample splitting, cross-fitting, out-of-bag prediction, and leave-one-out estimation. All of these techniques have the same final goal: nuisance parameters estimated to construct the influence function  $\hat{\Psi}_i$  for observation  $i$  [for the ATE case,  $\hat{\mu}(w, \mathbf{X}_i)$  and  $\hat{e}(\mathbf{X}_i)$ ] should be estimated without using outcome data about observation  $i$ . When random forests are used to estimate the nuisance parameters, this is straightforward, since out-of-bag predictions (standard in random forest statistical packages) provide the predictions obtained using trees that were constructed without using observation  $i$ . When other types of ML models are used to estimate the nuisance parameters, cross-fitting or sample splitting advocates splitting the data into folds and estimating the nuisance parameters separately on all data except a left-out fold, and then predicting the nuisance parameters in the left-out fold. When there are as many folds as observations, this is known as leave-one-out estimation.

Although these two issues are helpful in traditional small-data applications, when ML is used to estimate nuisance parameters (because there are many covariates), these issues become much more salient. First, overfitting is more of a concern, and in particular, a single observation  $i$  can have a strong effect on the predictions made for covariates  $\mathbf{X}_i$  when the model is very flexible. Cross-fitting can solve this problem. Second, we should expect that, with many covariates relative to the number of observations, accurate estimation of nuisance parameters is harder to achieve. Thus, orthogonalization makes estimation more robust to these errors.

### 6.3. Heterogeneous Treatment Effects

Another place where machine learning can be very useful is in uncovering treatment effect heterogeneity, where we focus on heterogeneity with respect to observable covariates. Examples of questions include: Which individuals benefit most from a treatment? For which individuals is the treatment effect positive? How do treatment effects change with covariates? Understanding treatment effect heterogeneity can be useful for basic scientific understanding or for estimating optimal policy assignments (for further discussion, see Athey & Imbens 2017b).

Continuing with the potential outcome notation from Section 6.1, we define the conditional ATE (CATE) as  $\tau(x) = E[\tau_i | \mathbf{X}_i = x]$ , where  $\tau_i = Y_i(1) - Y_i(0)$  is the treatment effect for individual  $i$ . The CATE is identified under the unconfoundedness assumption introduced in Section 6.1. Note that  $\tau_i$  cannot be observed for any unit; this “fundamental problem of causal inference” (Holland 1986, p. 947) is the source of an apparent difference between estimating heterogeneous treatment effects and predicting outcomes, which are typically observed for each unit.

We focus on three types of problems: (a) learning a low-dimensional representation of treatment effect heterogeneity and conducting hypothesis tests about this heterogeneity, (b) learning a flexible (nonparametric) estimate of  $\tau(x)$ , and (c) estimating an optimal policy allocating units to either treatment or control on the basis of covariates  $x$ .

An important issue in adapting ML methods to focus on causal parameters relates to the criterion function used in model selection. Predictive models typically use a mean squared error (MSE) criterion,  $\sum_i [Y_i - \hat{\mu}(\mathbf{X}_i)]^2 / N$ , to evaluate performance. Although the MSE in a held-out test set is a noisy estimate to the population expectation of the MSE in an independent set, the sample average MSE is a good, that is, unbiased, approximation that does not rely on further assumptions (beyond independence of observations), and the standard error of the squared errors in the test set accurately captures the uncertainty in the estimate. In contrast, consider the problem of estimating the CATE in observational studies. It would be natural to use as a criterion function the MSE of treatment effects,  $\sum_i [\tau_i - \hat{\tau}(\mathbf{X}_i)]^2 / N$ , where  $\hat{\tau}(x)$  is the estimate of the CATE. However, this criterion is infeasible, since we do not observe unit-level causal effects. Furthermore, there is no simple, model-free unbiased estimate of this criterion in observational studies. For this reason, comparing estimators and, as a result, developing regularization strategies are substantially harder challenges in settings where we are interested in structural or causal parameters than in settings where we are interested in predictive performance.

These difficulties in finding effective cross-validation strategies are not always insurmountable, but they lead to a need to carefully adapt and modify basic regularization methods to address the questions of interest. Athey & Imbens (2016) propose several different possible criteria to use for optimizing splits of the covariate space, as well as for cross-validation. A first insight is that, when conducting model selection, it is only necessary to compare models. The  $\tau_i^2$  term (which would be difficult to estimate) cancels out when comparing two estimators, say,  $\hat{\tau}'(x)$  and  $\hat{\tau}''(x)$ . The remaining terms are linear in  $\tau_i$ , and the expected value of  $\tau_i$  can be estimated. If we define

what Athey & Imbens (2016) call the transformed outcome,

$$Y_i^* = W_i \frac{Y_i}{e(\mathbf{X}_i)} - (1 - W_i) \frac{Y_i}{1 - e(\mathbf{X}_i)},$$

then  $\mathbb{E}[Y_i^*|\mathbf{X}_i] = \mathbb{E}[\tau_i|\mathbf{X}_i]$ . When the propensity score is unknown, it must be estimated, which implies that a criterion based on an estimate of the MSE of the CATE will depend on modeling choices.

Athey & Imbens (2016) build on this insight and propose several different estimators for the relative MSE of estimators for the CATE. They develop a method, which they call causal tree, for learning a low-dimensional representation of treatment effect heterogeneity, which provides reliable confidence intervals for the parameters that it estimates. Their paper builds on regression tree methods, creating a partition of the covariate space and then estimating treatment effects in each element of the partition. Unlike in regression trees optimized for prediction, the splitting rule optimizes for finding splits associated with treatment effect heterogeneity. In addition, the method relies on sample splitting; half of the data are used to estimate the tree structure, and the other half (the estimation sample) are used to estimate treatment effects in each leaf. The tree is pruned using cross-validation, just as in standard regression trees, but the criterion for evaluating the performance of the tree in held-out data is based on treatment effect heterogeneity, rather than predictive accuracy.

Some advantages of the causal tree method are similar to advantages of regression trees. These advantages are easy to explain; in the case of a randomized experiment, the estimate in each leaf is simply the sample ATE. A disadvantage is that the tree structure is somewhat arbitrary; there may be many partitions of the data that exhibit treatment effect heterogeneity, and taking a slightly different subsample of the data might lead to a different estimated partition. The approach of estimating simple models in the leaves of shallow trees can be applied to other types of models; Zeileis et al. (2008) provide an early version of this idea, although their paper does not provide theoretical guarantees or confidence intervals.

For some purposes, it is desirable to have a smooth estimate of  $\tau(x)$ . For example, if a treatment decision must be made for a particular individual with covariates  $x$ , a regression tree may give a biased estimate for that individual given that the individual may not be in the center of the leaf, and that the leaf may contain other units that are distant in covariate space. In the traditional econometrics literature, nonparametric estimation could be accomplished through kernel estimation or matching techniques. However, the theoretical and practical properties of these techniques are poor with many covariates. Wager & Athey (2017) introduce causal forests. Essentially, a causal forest is the average of a large number of causal trees, where trees differ from one another due to subsampling. Similar to prediction forests, a causal forest can be thought of as a version of a nearest neighbor matching method, but one where there is a data-driven approach to determine which dimensions of the covariate space are important to match on. Wager & Athey (2017) establish asymptotic normality of the estimator (so long as tree estimation is honest, making use of sample splitting for each tree) and provide an estimator for the variance of estimates so that confidence intervals can be constructed.

A challenge with forests is that it is difficult to describe the output, since the estimated CATE function  $\hat{\tau}(x)$  may be quite complex. However, in some cases, one might wish to test simpler hypotheses, such as the hypothesis that the top 10% of individuals ranked by their CATE have a different average CATE than the rest of the population. Chernozhukov et al. (2018b) provide methods for testing this type of hypothesis.

As described above in our presentation of regression forests, Athey et al. (2016b) extend the framework of causal forests to analyze nonparametric parameter heterogeneity in models where

the parameter of interest can be estimated by maximum likelihood or GMM. As an application, they highlight the case of instrumental variables. Friedberg et al. (2018) extend local linear regression forests to the problem of heterogeneous treatment effects, so that regularity in the function  $\tau(x)$  can be better exploited.

An alternative approach to estimating parameter heterogeneity in instrumental variables models was proposed by Hartford et al. (2016), who use an approach based on neural nets, although distributional theory is not available for that estimator. Other possible approaches to estimating conditional ATEs can be used when the structure of the heterogeneity is assumed to take a simple form. Targeted maximum likelihood (van der Laan & Rubin 2006) is one approach to this, while Imai & Ratkovic (2013) propose using LASSO to uncover heterogeneous treatment effects. Künzel et al. (2017) propose an ML approach using meta-learners. Another popular alternative that takes a Bayesian approach is Bayesian additive regression trees (BART), developed by Chipman et al. (2010) and applied to causal inference by Hill (2011) and Green & Kern (2012). A recent and promising method is the R-learner proposed by Nie & Wager (2019), which first estimates the two nuisance components, the conditional outcome mean and the propensity score, using flexible ML prediction methods and then focuses on a loss function that isolates the causal effects of interest from these nuisance components.

A main motivation for understanding treatment effect heterogeneity is that the CATE can be used to define policy assignment functions, that is, functions that map from the observable covariates of individuals to policy assignments. A simple way to define a policy is to estimate  $\hat{\tau}(x)$  and to assign the treatment to all individuals with positive values of  $\hat{\tau}(x)$ , where the estimate should be augmented with any costs of being in the treatment or control group. Hirano & Porter (2009) show that this is optimal under some conditions. A concern with this approach, depending on the method used to estimate  $\hat{\tau}(x)$ , is that the policy may be very complex and is not guaranteed to be smooth.

Kitagawa & Tetenov (2015) focus on estimating the optimal policy from a class of potential policies of limited complexity in an observational study with known propensity scores. The goal is to select a policy function to minimize the loss from failing to use the (infeasible) ideal policy, referred to as the regret of the policy. Athey & Wager (2017) also study policies with limited complexity; accommodate other constraints, such as budget constraints on the treatment; and propose an algorithm for estimating optimal policies. They provide bounds on the performance of their algorithm for the case where the data come from an observational study under confoundedness, and the propensity score is unknown. They also extend the analysis to settings that do not satisfy unconfoundedness, for example, settings where there is an instrumental variable.

Athey & Wager (2017) show how bringing in insights from semiparametric efficiency theory enables tighter bounds on performance than the ML literature, thus narrowing down substantially the set of algorithms that might achieve the regret bound. For the case of unconfoundedness, the policy estimation procedure recommended by Athey & Wager (2017) can be written as follows, where  $\Pi$  is the set of functions  $\pi : \mathbb{X} \rightarrow \{0, 1\}$ , and  $\hat{\Psi}_i$  is defined as above and makes use of cross-fitting as well as orthogonalization:

$$\max_{\pi \in \Pi} \sum_i [2\pi(\mathbf{X}_i) - 1] \cdot \hat{\Psi}_i. \quad 4.$$

The topic of optimal policy estimation has received some attention in the ML literature, focusing on data from observational studies with unconfoundedness, including those of Strehl et al. (2010), Dudik et al. (2011, 2014), Li et al. (2012, 2014), Swaminathan & Joachims (2015), Jiang & Li (2016), Thomas & Brunskill (2016), and Kallus (2017). Zhou et al. (2018) analyze the case with more than two treatment arms, extending the efficiency results of Athey & Wager (2017).

One insight that comes out of the ML approach to this problem is that the optimization problem in Equation 4 can be reframed as a classification problem and thus solved with off-the-shelf classification tools (for details, see Athey & Wager 2017).

## 7. EXPERIMENTAL DESIGN, REINFORCEMENT LEARNING, AND MULTI-ARMED BANDITS

ML methods have recently made substantial contributions to experimental design, with multi-armed bandits becoming more popular, especially in online experiments. Thompson sampling (Thompson 1933, Scott 2010) and upper confidence bounds (UCBs) (Lai & Robbins 1985) can be viewed as a simple example of reinforcement learning (Sutton & Barto 1998) where successful assignment decisions are rewarded by sending more units to the corresponding treatment arm.

### 7.1. A/B Testing Versus Multi-Armed Bandits

Traditionally, much experimentation is done by assigning a predetermined number of units to each of a number of treatment arms. There would often be just two treatment arms. After the outcomes are measured, the average effect of the treatment would be estimated using the difference in average outcomes by treatment arm. This is a potentially very inefficient way of experimentation, where we waste units by assigning them to treatment arms that we already know with a high degree of confidence to be inferior to some of the other arms. Modern methods for online experimentation focus on balancing exploration of new treatments with exploitation of treatments currently assessed to be of high quality. Suppose that what we are interested in is primarily finding a treatment that is good among the set of treatments considered, rather than in estimation of expected outcomes for the full set of treatments. Moreover, suppose that we measure the outcomes quickly after the treatments have been assigned, and suppose that the units arrive sequentially for assignment to a treatment. After outcomes for half the units have been observed, we may have a pretty good idea which of the treatments are still candidates for the optimal treatment. Exposing more units to treatments that are no longer competitive is suboptimal for both exploration and exploitation purposes: It does not help us distinguish between the remaining candidate optimal treatments, and it exposes those units to inferior treatments.

Multi-armed bandit approaches (Thompson 1933, Scott 2010) attempt to improve over this static design. In the extreme case, the assignment for each unit depends on all of the information learned up to that point. Given that information, and given a parametric model for the outcomes for each treatment and a prior for the parameters of these models, we can estimate the probability of each treatment being the optimal one. Thompson sampling suggests assigning the next unit to each treatment with probability equal to the probability that that particular treatment is the optimal one. This means that the probability of assignment to a treatment arm that we are confident is inferior to some of the other treatments is low, and eventually, all new units will be assigned to the optimal treatment with probability close to one.

To provide some more intuition, consider a case with  $K$  treatments where the outcome is binary, so the model is a binomial distribution with treatment-arm-specific success probability  $p_k$ , for  $k = 1, \dots, K$ . If the prior distribution for all probabilities is uniform, then the posterior distribution for the success probability of arm  $k$ , given  $M_k$  successes in the  $N_k$  trials concluded so far, is Beta with parameters  $M_k + 1$  and  $N_k - M_k + 1$ . Given that the Beta distribution is simple to approximate by simulation, the probability that treatment arm  $k$  is the optimal one (the one with the highest success probability) is  $\text{pr}(p_k = \max_{m=1}^K p_m)$ .

We can simplify the calculations by updating the assignment probabilities only after seeing several new observations. That is, we reevaluate the assignment probabilities after a batch of new

observations has come in, all based on the same assignment probabilities. From this perspective, we can view a standard A/B experiment as one where the batch is the full set of observations. This makes it clear that to, at least occasionally, update the assignment probabilities to avoid sending units to inferior treatments is a superior strategy.

An alternative approach is to use the UCB (Lai & Robbins 1985) approach. In this case, we construct a  $100(1 - p)\%$  confidence interval for the population average outcome  $\mu_k$  for each treatment arm. We then collect the upper bounds of these confidence intervals for each treatment arm and assign the next unit to the treatment arm with the highest value for the UCBs. As we get more and more data, we let one minus the level of the confidence intervals  $p$  go to zero slowly. With UCB methods, we need to be more careful if we wish to update assignments only after batches of units have come in. If two treatment arms have very similar UCBs, assigning a large number of units to the one that has a slightly higher UCB may not be satisfactory: In this case, Thompson sampling would assign similar numbers of units to both of these treatment arms. More generally, the stochastic nature of the assignment under Thompson sampling, compared to the deterministic assignment in the UCB approach, has conceptual advantages for the ability to do randomization inference (e.g., Athey & Imbens 2017a).

## 7.2. Contextual Bandits

The most important extension of multi-armed bandits is to settings where we observe features of the units that can be used in the assignment mechanism. If treatment effects are heterogeneous, and if that heterogeneity is associated with observed characteristics of the units, then there may be substantial gains from assigning units to different treatments based on these characteristics (for details, see Dimakopoulou et al. 2018).

A simple way to incorporate covariates would be to build a parametric model for the expected outcomes in each treatment arm (the reward function), estimate that given the current data, and infer from there the probability that a particular arm is optimal for a new unit conditional on the characteristics of that unit. This is conceptually a straightforward way to incorporate characteristics, but it has some drawbacks. The main concern is that such methods may implicitly rely substantially on the model being correctly specified. It may be the case that the data for one treatment arm come in with a particular distribution of the characteristics, but they are used to predict outcomes for units with very different characteristics (for discussion, see Bastani & Bayati 2015). A risk is that, if the algorithm estimates a simple linear model mapping characteristics to outcomes, then the algorithm may suggest a great deal of certainty about outcomes for an arm in a region of characteristic space where that treatment arm has never been observed. This can lead the algorithm to never experiment with the arm in that region, allowing for the possibility that the algorithm never corrects its mistake and fails to learn the true optimal policy even in large samples.

As a result, one should be careful in building a flexible model relating the characteristics to the outcomes. Dimakopoulou et al. (2017) highlight the benefits of using random forests as a way to avoid making functional form assumptions.

Beyond this issue, several novel considerations arise in contextual bandits. Because the assignment rules as a function of the features change as more units arrive and tend to assign more units to a given arm in regions of the covariate space where the assignment rule has performed well in the past, particular care has to be taken to eliminate biases in the estimation of the reward function. Thus, although there is formal randomization, the issues concerning robust estimation of conditional average causal effects in observational studies become relevant in this case. One solution, motivated by the literature on causal inference, is to use propensity score weighting of outcome models. Dimakopoulou et al. (2018) study bounds on the performance of contextual bandits

using doubly robust estimation (propensity-weighted outcome modeling) and also demonstrate on a number of real-world data sets that propensity weighting improves performance.

Another insight is that it can be useful to make use of simple assignment rules, particularly in early stages of bandits, because complex assignment rules can lead to confounding later. In particular, if a covariate is related to outcomes and is used in assignment, then later estimation must control for this covariate to eliminate bias. For this reason, LASSO, which selects a sparse model, can perform better than ridge, which places weights on more covariates, when estimating an outcome model that will be used to determine the assignment of units in subsequent batches. Finally, flexible outcome models can be important in certain settings; random forests can be a good alternative in these cases.

## 8. MATRIX COMPLETION AND RECOMMENDER SYSTEMS

The methods that we discuss above are primarily for settings where we observe information on several units in the form of a single outcome and a set of covariates or features, what is known in the econometrics literature as a cross-section setting. There are also many interesting new methods for settings that resemble what are, in the econometric literature, referred to as longitudinal or panel data settings. In this section, we discuss a canonical version of that problem and consider some specific methods.

### 8.1. The Netflix Problem

The Netflix Prize competition was set up in 2006 (Bennett & Lanning 2007) and asked researchers to use a training data set to develop an algorithm that improved on the Netflix algorithm for recommending movies by providing predictions for movie ratings. Researchers were given a training data set that contained movie and individual characteristics, as well as movie ratings, and were asked to predict ratings for movie–individual pairs for which they were not given the ratings. Because of the magnitude of the prize, \$1,000,000, this competition and the associated problem generated a lot of attention, and the development of new methods for this type of setting accelerated substantially as a result. The winning solutions, and those that were competitive with the winners, had some key features. First, they relied heavily on model averaging. Second, many of the models included matrix factorization and nearest neighbor methods.

Although it may appear at first to be a problem that is very distinct from the type of problem studied in econometrics, one can cast many econometric panel data in a similar form. In settings where researchers are interested in causal effects of a binary treatment, one can think of the realized data as consisting of two incomplete potential outcome matrices, one for the outcomes given the treatment and one for the outcomes given the control treatment. Thus, the problem of estimating the ATEs can be cast as a matrix completion problem. Suppose that we observe outcomes on  $N$  units over  $T$  time periods, with the outcome for unit  $i$  at time period  $t$  denoted by  $Y_{it}$ , and a binary treatment, denoted by  $W_{it}$ , with

$$\mathbf{Y} = \begin{pmatrix} Y_{11} & Y_{12} & Y_{13} & \dots & Y_{1T} \\ Y_{21} & Y_{22} & Y_{23} & \dots & Y_{2T} \\ Y_{31} & Y_{32} & Y_{33} & \dots & Y_{3T} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ Y_{N1} & Y_{N2} & Y_{N3} & \dots & Y_{NT} \end{pmatrix} \quad \mathbf{W} = \begin{pmatrix} 1 & 1 & 0 & \dots & 1 \\ 0 & 0 & 1 & \dots & 0 \\ 1 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & 1 & \dots & 0 \end{pmatrix}.$$

We can think of there being two matrices with potential outcomes,

$$\mathbf{Y}(0) = \begin{pmatrix} ? & ? & Y_{13} & \dots & ? \\ Y_{21} & Y_{22} & ? & \dots & Y_{2T} \\ ? & Y_{32} & ? & \dots & Y_{3T} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ ? & Y_{N2} & ? & \dots & Y_{NT} \end{pmatrix} \quad (\text{potential control outcome})$$

and

$$\mathbf{Y}(1) = \begin{pmatrix} Y_{11} & Y_{12} & ? & \dots & Y_{1T} \\ ? & ? & Y_{23} & \dots & ? \\ Y_{31} & ? & Y_{33} & \dots & ? \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ Y_{N1} & ? & Y_{N3} & \dots & ? \end{pmatrix} \quad (\text{potential treated outcome}).$$

Thus, the problem of estimating causal effects becomes one of imputing missing values in a matrix.

The ML literature has developed effective methods for matrix completion in settings with both  $N$  and  $T$  large and a large fraction of missing data. We discuss some of these methods in the next section, as well as their relation to the econometrics literature.

## 8.2. Matrix Completion Methods for Panel Data

The matrix completion literature has focused on using low-rank representations for the complete data matrix. Let us consider the case without covariates, that is, no characteristics of the units or time periods. Let  $\mathbf{L}$  be the matrix of expected values, and  $\mathbf{Y}$  be the observed data matrix. The observed values are assumed to be equal to the corresponding values of the complete data matrix, possibly with error:

$$Y_{it} = \begin{cases} L_{it} + \varepsilon_{it} & \text{if } W_{it} = 1, \\ 0 & \text{otherwise.} \end{cases}$$

Using the singular value decomposition,  $\mathbf{L} = \mathbf{U}\mathbf{S}\mathbf{V}^\top$ , where  $\mathbf{U}$  is an  $N \times N$  matrix,  $\mathbf{V}$  is a  $T \times T$  matrix, and  $\mathbf{S}$  is an  $N \times T$  matrix with rank  $R$ , with the only nonzero elements on the diagonal (the singular values). We are not interested in estimating the matrices  $\mathbf{U}$  and  $\mathbf{V}$ , only in the product  $\mathbf{U}\mathbf{S}\mathbf{V}^\top$ , and possibly in the singular values, the diagonal elements of  $\mathbf{S}$ . Obviously some regularization is required, and an effective one is to use the nuclear norm  $\|\cdot\|_*$ , which is equal to the sum of the singular values. Building on the ML literature, Candès & Recht (2009); Mazumder et al. (2010), and Athey et al. (2017a) focus on estimating  $\mathbf{L}$  by minimizing

$$\min_{\mathbf{L}} \left\{ \sum_{(i,t) \in \mathcal{O}} (Y_{it} - L_{it})^2 + \lambda \|\mathbf{L}\|_* \right\},$$



where  $\lambda$  is a penalty parameter chosen through cross-validation. Using the nuclear norm in this case, rather than the rank of the matrix  $\mathbf{L}$ , is important for computational reasons. Using the Frobenius norm equal to the sum of the squares of the singular values would not work because it is equal to the sum of the squared values of the matrix and thus would lead to imputing all missing values as zeros. For the nuclear norm case, there are effective algorithms that can deal with large  $N$  and large  $T$  (see Candès & Recht 2009, Mazumder et al. 2010).

### 8.3. The Econometrics Literature on Panel Data and Synthetic Control Methods

The econometrics literature has studied these problems from a number of different perspectives. The panel data literature has traditionally focused on fixed-effect methods and has generalized these to models with multiple latent factors (Bai & Ng 2002, 2017; Bai 2003) that are essentially the same as the low-rank factorizations in the ML literature. The difference is that in the econometrics literature there has been more focus on actually estimating the factors and using normalizations that allow for their identification. It is typically assumed that there is a fixed number of factors.

The synthetic control literature has studied similar settings but focused on the case with only missing values for a single row of the matrix  $\mathbf{Y}$ . Abadie et al. (2010, 2015) propose imputing these using a weighted average of the outcomes for other units in the same period. Doudchenko & Imbens (2016) show that the Abadie et al. (2015) methods can be viewed as regressing the outcomes for the last row on outcomes for the other units and using the regression estimates from that to impute the missing values, in what Athey et al. (2017a) call the vertical regression. This contrasts with a horizontal regression, common in the program evaluation literature, where outcomes in the last period are regressed on outcomes in earlier periods, and those estimates are used to impute the missing values. In contrast to both the horizontal and vertical regression approaches, the matrix completion approach, in principle, attempts to exploit both stable patterns over time and stable patterns between units in imputing the missing values and can also deal directly with more complex missing data patterns.

### 8.4. Demand Estimation in Panel Data

A large literature in economics and marketing focuses on estimating consumer preferences using data about their choices. A typical paper analyzes the discrete choice of a consumer who selects a single product from a set of prespecified imperfect substitutes, e.g., laundry detergent, personal computers, or cars (for a review, see, e.g., Keane 2013). The literature typically focuses on one product category at a time and models choices among a small number of products. This literature often focuses on estimating cross-price elasticities, so that counterfactuals about firm mergers or price changes can be analyzed. Although it is common to incorporate individual-specific preferences for observable characteristics, such as prices and other product characteristics, there are typically a small number of latent variables in the models. A standard set-up starts with consumer  $i$ 's utility for product  $j$  at time  $t$ , where

$$U_{ijt} = \mu_{ij} - \phi_i p_{jt} + \epsilon_{ijt},$$

where  $\epsilon_{ijt}$  has an extreme value distribution and is independently and identically distributed across consumers, products, and time. The term  $\mu_{ij}$  is consumer  $i$ 's mean utility for product  $j$ ,  $\phi_i$  is consumer  $i$ 's price sensitivity, and  $p_{jt}$  is the price of product  $j$  at time  $t$ . If the consumer selects the

item with highest utility, then

$$Pr(Y_{ijt} = j) = \frac{\exp^{\mu_{ij} - \phi_i p_{jt}}}{\sum_{j'} \exp^{\mu_{ij'} - \phi_i p_{jt}}}.$$

From the ML perspective, a panel data set with consumer choices might be studied using techniques from matrix completion, as described above. The model would draw inferences from products that had similar purchase patterns across consumers, as well as consumers who had similar purchase patterns across products. However, such models would typically not be well-suited to analyze the extent to which two products are substitutes or to analyze counterfactuals.

For example, Jacobs et al. (2014) propose using a related latent factorization approach to flexibly model consumer heterogeneity in the context of online shopping with a large assortment of products. They use data from a medium-sized online retailer. They consider 3,226 products and aggregate up to the category  $\times$  brand level to reduce to 440 products. They do not model responses to price changes or substitution between similar products; instead, in the spirit of the ML literature, they evaluate performance in terms of predicting which new products a customer will buy.

In contrast to this off-the-shelf application of ML to product choice, a recent literature has emerged that attempts to combine ML methods with insights from the economics literature on consumer choice, typically in panel data settings. A theme of this literature is that models that take advantage of some of the structure of the problem will outperform models that do not. For example, the functional form implied by the consumer choice model from economics places a lot of structure on how products within a category interact with one another. An increase in the price of one product affects other products in a particular way, implied by the functional form. To the extent that the restrictions implied by the functional form are good approximations to reality, they can greatly improve the efficiency of estimation. Incorporating the functional forms that have been established to be effective across decades of economic research can improve performance.

However, economic models have typically failed to incorporate all of the information that is available in a panel data set, the type of information that matrix completion methods typically exploit. In addition, computational issues have prevented economists from studying consumer choices across multiple product categories, even though, in practice, data about a consumer's purchases in one category are informative about the consumer's purchases in other categories; furthermore, the data can also reveal which products tend to have similar purchase patterns. Thus, the best-performing models from this new hybrid literature tend to exploit techniques from the matrix completion literature, in particular, matrix factorization.

To see how matrix factorization can augment a standard consumer choice model, we can write the utility of consumer  $i$  for product  $j$  at time  $t$  as

$$U_{ijt} = \beta_i' \theta_j - \rho_i' \alpha_j p_{jt} + \epsilon_{ijt},$$

where  $\beta_i$ ,  $\theta_j$ ,  $\rho_i$ , and  $\alpha_j$  are vectors of latent variables. The vector  $\theta_j$ , for example, can be interpreted as a vector of latent product characteristics for product  $j$ , while  $\beta_i$  represents consumer  $i$ 's latent preferences for those characteristics. The basic functional form for choice probabilities is unchanged, except that the utilities are now functions of the latent characteristics.

Such models had not been studied in the ML literature until recently, in part because the functional form for choice probabilities, which is nonlinear in a large number of latent parameters, makes computation challenging. In contrast, traditional ML models might treat all products as independently chosen (e.g., Gopalan et al. 2015), making computation much easier. Ruiz et al. (2017) apply state-of-the-art computational techniques from ML (in particular, stochastic

gradient descent and variational inference) together with several approximations to make the method scalable to thousands of consumers making choices over thousands of items in dozens or hundreds of shopping trips per consumer. Ruiz et al. (2017) do not make use of any data about the categories of products; they attempt to learn from the data (which incorporate substantial price variation) which products are substitutes or complements. In contrast, Athey et al. (2017b) incorporate information about product categories and impose the assumption that consumers buy only one product per category on a given trip; they also introduce a nested logit structure, which allows utilities to be correlated across products within a category, thus better accounting for consumers' choices about whether to purchase a category at all.

A closely related approach is taken by Wan et al. (2017). They use a latent factorization approach that incorporates price variation. They model consumer choice as a three-stage process: (a) Choose whether to buy from the category, (b) choose which item in the category, and (c) choose the number of the item to purchase. The paper uses customer loyalty transaction data from two different data sets. In all of these approaches, using the utility maximization approach from economics makes it possible to perform traditional analyses such as analyzing the impact of price changes on consumer welfare. A complementary approach to one based on latent product characteristics is the work by Semenova et al. (2018), who consider observational high-dimensional product attributes (e.g., text descriptions and images) rather than latent features.

## 9. TEXT ANALYSIS

There is a large ML literature on analyzing text data. It is beyond the scope of this article to fully describe this literature; Gentzkow et al. (2017) provide an excellent recent review. In this section, we provide a high-level overview.

To start, we consider a data set consisting of  $N$  documents, indexed by  $i = 1, \dots, N$ . Each document contains a set of words. One way to represent the data is as a  $N \times T$  matrix, denoted  $C$ , where  $T$  is the number of words in the language, and where each element of the matrix is an indicator for whether word  $t$  appears in document  $i$ . Such a representation would lose information by ignoring the ordering of the words in the text. Richer representations might let  $T$  be the number of bigrams, where a bigram is a pair of words that appear adjacent to one another in the document, or sequences of three or more words.

There are two types of exercise we can do with this type of data. One is unsupervised learning, and the other is supervised. For the unsupervised case, the goal would be to find a lower-rank representation of the matrix  $C$ . Given that a low-rank matrix can be well approximated by a factor structure, as discussed above, this is equivalent to finding a set of  $k$  latent characteristics of documents (denoted  $\beta$ ) and a set of latent weights on these topics (denoted  $\theta$ ) such that the probability that word  $t$  appears in document  $i$  is a function of  $\theta_i' \beta_j$ . This view of the problem basically turns it into a matrix completion problem; we would say that a particular representation performs well if we hold out a test set of randomly selected elements of  $C$ , and the model predicts well those held-out elements. All of the methods described above for matrix completion can be applied in this case.

One implementation of these ideas is referred to as a topic model (for a review, see Blei & Lafferty 2009). This model specifies a particular generative model of the data. In the model, there are several topics, which are latent variables. Each topic is associated with a distribution of words. An article is characterized by weights on each topic. The goal of a topic model is to estimate the latent topics, the distribution over words for each topic, and the weights for each article. A popular model that does this is known as the latent Dirichlet allocation model.

More recently, more complex models of language have emerged, following the theme that, although simple ML models perform quite well, incorporating problem-specific structure is often helpful and is typically done in state-of-the art ML in popular application areas. Broadly, these are known as word embedding methods. These attempt to capture latent semantic structure in language (see Mnih & Hinton 2007; Mnih & Teh 2012; Mikolov et al. 2013a,b,c; Mnih & Kavukcuoglu 2013; Levy & Goldberg 2014; Pennington et al. 2014; Vilnis & McCallum 2015; Arora et al. 2016; Barkan 2016; Bamler & Mandt 2017). Consider the neural probabilistic language model of Bengio et al. (2003, 2006). This model specifies a joint probability of sequences of words, parameterized by a vector representation of the vocabulary. Vector representations of words (also known as distributed representations) can incorporate ideas about word usage and meaning (Harris 1954, Firth 1957, Bengio et al. 2003, Mikolov et al. 2013b).

Another class of models uses supervised learning methods. These methods are used when there is a specific characteristic that the researcher would like to learn from the text. Examples might include favorability of a review, political polarization of text spoken by legislators, or whether a tweet about a company is positive or negative. Then, the outcome variable is a label that contains the characteristic of interest. A simple supervised learning model takes the data matrix  $C$ , views each document  $i$  as a unit of observation, and treats the columns of  $C$  (each corresponding to indicators for whether a particular word is in a document) as the covariates in the regression. Since  $T$  is usually much greater than  $N$ , it is important to use ML methods that allow for regularization. Sometimes, other types of dimension reduction techniques are used in advance of applying a supervised learning method (e.g., unsupervised topic modeling).

Another approach is to think of a generative model, where we think of the words in the document as a vector of outcomes, and where the characteristics of interest about the document determine the distribution of words, as in the topic model literature. An example of this approach is the supervised topic model, where information about the observed characteristics in a training data set are incorporated into the estimation of the generative model. The estimated model can then be used to predict those characteristics in a test data set of unlabeled documents (for more details, see Blei & Lafferty 2009).

## 10. CONCLUSION

There is a fast-growing ML literature that has much to offer empirical researchers in economics. In this review, we describe some of the methods that we view as most useful for economists, and that we view as important to include in the core graduate econometrics sequences. Being familiar with these methods will allow researchers to do more sophisticated empirical work and to communicate more effectively with researchers in other fields.

## DISCLOSURE STATEMENT

The authors are not aware of any affiliations, memberships, funding, or financial holdings that might be perceived as affecting the objectivity of this review.

## ACKNOWLEDGMENTS

We are grateful to Sylvia Klosin for comments. This research was generously supported by Office of Naval Research grant N00014-17-1-2131 to Sylvia Klosin and Michael Pollmann and by the Sloan Foundation.

## LITERATURE CITED

- Abadie A, Cattaneo MD. 2018. Econometric methods for program evaluation. *Annu. Rev. Econ.* 10:465–503
- Abadie A, Diamond A, Hainmueller J. 2010. Synthetic control methods for comparative case studies: estimating the effect of California's tobacco control program. *J. Am. Stat. Assoc.* 105:493–505
- Abadie A, Diamond A, Hainmueller J. 2015. Comparative politics and the synthetic control method. *Am. J. Political Sci.* 59:495–510
- Abadie A, Imbens GW. 2011. Bias-corrected matching estimators for average treatment effects. *J. Bus. Econ. Stat.* 29:1–11
- Alpaydin E. 2009. *Introduction to Machine Learning*. Cambridge, MA: MIT Press
- Angrist JD, Pischke JS. 2008. *Mostly Harmless Econometrics: An Empiricist's Companion*. Princeton, NJ: Princeton Univ. Press
- Arjovsky M, Bottou L. 2017. Towards principled methods for training generative adversarial networks. arXiv:1701.04862 [stat.ML]
- Arora S, Li Y, Liang Y, Ma T. 2016. RAND-WALK: a latent variable model approach to word embeddings. *Trans. Assoc. Comput. Linguist.* 4:385–99
- Athey S. 2017. Beyond prediction: using big data for policy problems. *Science* 355:483–85
- Athey S. 2019. The impact of machine learning on economics. In *The Economics of Artificial Intelligence: An Agenda*, ed. AK Agrawal, J Gans, A Goldfarb. Chicago: Univ. Chicago Press. In press
- Athey S, Bayati M, Doudchenko N, Imbens G, Khosravi K. 2017a. Matrix completion methods for causal panel data models. arXiv:1710.10251 [math.ST]
- Athey S, Bayati M, Imbens G, Zhaonan Q. 2019. *Ensemble methods for causal effects in panel data settings*. NBER Work. Pap. 25675
- Athey S, Blei D, Donnelly R, Ruiz F. 2017b. Counterfactual inference for consumer choice across many product categories. *AEA Pap. Proc.* 108:64–67
- Athey S, Imbens G. 2016. Recursive partitioning for heterogeneous causal effects. *PNAS* 113:7353–60
- Athey S, Imbens G, Wager S. 2016a. Efficient inference of average treatment effects in high dimensions via approximate residual balancing. arXiv:1604.07125 [math.ST]
- Athey S, Imbens GW. 2017a. The econometrics of randomized experiments. In *Handbook of Economic Field Experiments*, Vol. 1, ed. E Duflo, A Banerjee, pp. 73–140. Amsterdam: Elsevier
- Athey S, Imbens GW. 2017b. The state of applied econometrics: causality and policy evaluation. *J. Econ. Perspect.* 31:3–32
- Athey S, Mobius MM, Pál J. 2017c. *The impact of aggregators on internet news consumption*. Unpublished manuscript, Grad. School Bus., Stanford Univ., Stanford, CA
- Athey S, Tibshirani J, Wager S. 2016b. Generalized random forests. arXiv:1610.01271 [stat.ME]
- Athey S, Wager S. 2017. Efficient policy learning. arXiv:1702.02896 [math.ST]
- Bai J. 2003. Inferential theory for factor models of large dimensions. *Econometrica* 71:135–71
- Bai J, Ng S. 2002. Determining the number of factors in approximate factor models. *Econometrica* 70:191–221
- Bai J, Ng S. 2017. Principal components and regularized estimation of factor models. arXiv:1708.08137 [stat.ME]
- Bamler R, Mandt S. 2017. Dynamic word embeddings via skip-gram filtering. In *Proceedings of the 34th International Conference on Machine Learning*, pp. 380–89. La Jolla, CA: Int. Mach. Learn. Soc.
- Barkan O. 2016. Bayesian neural word embedding. arXiv:1603.06571 [math.ST]
- Bastani H, Bayati M. 2015. *Online decision-making with high-dimensional covariates*. Work. Pap., Univ. Penn./Stanford Grad. School Bus., Philadelphia/Stanford, CA
- Bell RM, Koren Y. 2007. Lessons from the Netflix prize challenge. *ACM SIGKDD Explor. Newsl.* 9:75–79
- Belloni A, Chernozhukov V, Hansen C. 2014. High-dimensional methods and inference on structural and treatment effects. *J. Econ. Perspect.* 28:29–50
- Bengio Y, Ducharme R, Vincent P, Janvin C. 2003. A neural probabilistic language model. *J. Mach. Learn. Res.* 3:1137–55
- Bengio Y, Schwenk H, Senécal JS, Morin F, Gauvain JL. 2006. Neural probabilistic language models. In *Innovations in Machine Learning: Theory and Applications*, ed. DE Holmes, pp. 137–86. Berlin: Springer

- Bennett J, Lanning S. 2007. The Netflix prize. In *Proceedings of KDD Cup and Workshop 2007*, p. 35. New York: ACM
- Bertsimas D, King A, Mazumder R. 2016. Best subset selection via a modern optimization lens. *Ann. Stat.* 44:813–52
- Bickel P, Klaassen C, Ritov Y, Wellner J. 1998. *Efficient and Adaptive Estimation for Semiparametric Models*. Berlin: Springer
- Bierens HJ. 1987. Kernel estimators of regression functions. In *Advances in Econometrics: Fifth World Congress*, Vol. 1, ed. TF Bewley, pp. 99–144. Cambridge, UK: Cambridge Univ. Press
- Blei DM, Lafferty JD. 2009. Topic models. In *Text Mining: Classification, Clustering, and Applications*, ed. A Srivastava, M Sahami, pp. 101–24. Boca Raton, FL: CRC Press
- Bottou L. 1998. Online learning and stochastic approximations. In *On-Line Learning in Neural Networks*, ed. D Saad, pp. 9–42. New York: ACM
- Bottou L. 2012. Stochastic gradient descent tricks. In *Neural Networks: Tricks of the Trade*, ed. G Montavon, G Orr, K-R Müller, pp. 421–36. Berlin: Springer
- Breiman L. 1993. *Better subset selection using the non-negative garotte*. Tech. Rep., Univ. Calif., Berkeley
- Breiman L. 1996. Bagging predictors. *Mach. Learn.* 24:123–40
- Breiman L. 2001a. Random forests. *Mach. Learn.* 45:5–32
- Breiman L. 2001b. Statistical modeling: the two cultures (with comments and a rejoinder by the author). *Stat. Sci.* 16:199–231
- Breiman L, Friedman J, Stone CJ, Olshen RA. 1984. *Classification and Regression Trees*. Boca Raton, FL: CRC Press
- Burkov A. 2019. *The Hundred-Page Machine Learning Book*. Quebec City, Can.: Andriy Burkov
- Candès E, Tao T. 2007. The Dantzig selector: statistical estimation when  $p$  is much larger than  $n$ . *Ann. Stat.* 35:2313–51
- Candès EJ, Recht B. 2009. Exact matrix completion via convex optimization. *Found. Comput. Math.* 9:717
- Chamberlain G. 2000. Econometrics and decision theory. *J. Econom.* 95:255–83
- Chen X. 2007. Large sample sieve estimation of semi-nonparametric models. In *Handbook of Econometrics*, Vol. 6B, ed. JJ Heckman, EE Learner, pp. 5549–632. Amsterdam: Elsevier
- Chernozhukov V, Chetverikov D, Demirer M, Duflo E, Hansen C, et al. 2016a. *Double machine learning for treatment and causal parameters*. Tech. Rep., Cent. Microdata Methods Pract., Inst. Fiscal Stud., London
- Chernozhukov V, Chetverikov D, Demirer M, Duflo E, Hansen C, et al. 2018a. Double/debiased machine learning for treatment and structural parameters. *Econom. J.* 21:C1–68
- Chernozhukov V, Chetverikov D, Demirer M, Duflo E, Hansen C, Newey W. 2017. Double/debiased/ Neyman machine learning of treatment effects. *Am. Econ. Rev.* 107:261–65
- Chernozhukov V, Demirer M, Duflo E, Fernandez-Val I. 2018b. *Generic machine learning inference on heterogeneous treatment effects in randomized experiments*. NBER Work. Pap. 24678
- Chernozhukov V, Escanciano JC, Ichimura H, Newey WK. 2016b. Locally robust semiparametric estimation. arXiv:1608.00033 [math.ST]
- Chernozhukov V, Newey W, Robins J. 2018c. Double/de-biased machine learning using regularized Riesz representer. arXiv:1802.08667 [stat.ML]
- Chipman HA, George EI, McCulloch RE. 2010. Bart: Bayesian additive regression trees. *Ann. Appl. Stat.* 4:266–98
- Cortes C, Vapnik V. 1995. Support-vector networks. *Mach. Learn.* 20:273–97
- Dietterich TG. 2000. Ensemble methods in machine learning. In *Multiple Classifier Systems: First International Workshop, Cagliari, Italy, June 21–23*, pp. 1–15. Berlin: Springer
- Dimakopoulou M, Athey S, Imbens G. 2017. Estimation considerations in contextual bandits. arXiv:1711.07077 [stat.ML]
- Dimakopoulou M, Zhou Z, Athey S, Imbens G. 2018. Balanced linear contextual bandits. arXiv:1812.06227 [cs.LG]
- Doudchenko N, Imbens GW. 2016. *Balancing, regression, difference-in-differences and synthetic control methods: a synthesis*. NBER Work. Pap. 22791

- Dudik M, Erhan D, Langford J, Li L. 2014. Doubly robust policy evaluation and optimization. *Stat. Sci.* 29:485–511
- Dudik M, Langford J, Li L. 2011. Doubly robust policy evaluation and learning. In *Proceedings of the 28th International Conference on Machine Learning*, pp. 1097–104. La Jolla, CA: Int. Mach. Learn. Soc.
- Efron B, Hastie T. 2016. *Computer Age Statistical Inference*, Vol. 5. Cambridge, UK: Cambridge Univ. Press
- Efron B, Hastie T, Johnstone I, Tibshirani R. 2004. Least angle regression. *Ann. Stat.* 32:407–99
- Farrell MH, Liang T, Misra S. 2018. Deep neural networks for estimation and inference: application to causal effects and other semiparametric estimands. arXiv:1809.09953 [econ.EM]
- Firth JR. 1957. A synopsis of linguistic theory 1930–1955. In *Studies in Linguistic Analysis (Special Volume of the Philological Society)*, ed. JR Firth, pp. 1–32. Oxford, UK: Blackwell
- Friedberg R, Tibshirani J, Athey S, Wager S. 2018. Local linear forests. arXiv:1807.11408 [stat.ML]
- Friedman JH. 2002. Stochastic gradient boosting. *Comput. Stat. Data Anal.* 38:367–78
- Gentzkow M, Kelly BT, Taddy M. 2017. *Text as data*. NBER Work. Pap. 23276
- Goodfellow I, Pouget-Abadie J, Mirza M, Xu B, Warde-Farley D, et al. 2014. Generative adversarial nets. In *Advances in Neural Information Processing Systems*, Vol. 27, ed. Z Ghahramani, M Welling, C Cortes, ND Lawrence, KQ Weinberger, pp. 2672–80. San Diego, CA: Neural Inf. Process. Syst. Found.
- Gopalan P, Hofman J, Blei DM. 2015. Scalable recommendation with hierarchical Poisson factorization. In *Proceedings of the 31st Conference on Uncertainty in Artificial Intelligence, Amsterdam, Neth., July 12–16*, art. 208. Amsterdam: Assoc. Uncertain. Artif. Intell.
- Green DP, Kern HL. 2012. Modeling heterogeneous treatment effects in survey experiments with Bayesian additive regression trees. *Public Opin. Q.* 76:491–511
- Greene WH. 2000. *Econometric Analysis*. Upper Saddle River, NJ: Prentice Hall. 4th ed.
- Harris ZS. 1954. Distributional structure. *Word* 10:146–62
- Hartford J, Lewis G, Taddy M. 2016. Counterfactual prediction with deep instrumental variables networks. arXiv:1612.09596 [stat.AP]
- Hartigan JA, Wong MA. 1979. Algorithm as 136: a k-means clustering algorithm. *J. R. Stat. Soc. Ser. C* 28:100–8
- Hastie T, Tibshirani R, Friedman J. 2009. *The Elements of Statistical Learning*. Berlin: Springer
- Hastie T, Tibshirani R, Tibshirani RJ. 2017. Extended comparisons of best subset selection, forward stepwise selection, and the lasso. arXiv:1707.08692 [stat.ME]
- Hastie T, Tibshirani R, Wainwright M. 2015. *Statistical Learning with Sparsity: The Lasso and Generalizations*. New York: CRC Press
- Hill JL. 2011. Bayesian nonparametric modeling for causal inference. *J. Comput. Graph. Stat.* 20:217–40
- Hirano K, Porter JR. 2009. Asymptotics for statistical treatment rules. *Econometrica* 77:1683–701
- Hoerl AE, Kennard RW. 1970. Ridge regression: biased estimation for nonorthogonal problems. *Technometrics* 12:55–67
- Holland PW. 1986. Statistics and causal inference. *J. Am. Stat. Assoc.* 81:945–60
- Hornik K, Stinchcombe M, White H. 1989. Multilayer feedforward networks are universal approximators. *Neural Netw.* 2:359–66
- Imai K, Ratkovic M. 2013. Estimating treatment effect heterogeneity in randomized program evaluation. *Ann. Appl. Stat.* 7:443–70
- Imbens G, Wooldridge J. 2009. Recent developments in the econometrics of program evaluation. *J. Econ. Lit.* 47:5–86
- Imbens GW, Lemieux T. 2008. Regression discontinuity designs: a guide to practice. *J. Econom.* 142:615–35
- Imbens GW, Rubin DB. 2015. *Causal Inference in Statistics, Social, and Biomedical Sciences*. Cambridge, UK: Cambridge Univ. Press
- Jacobs B, Donkers B, Fok D. 2014. *Product Recommendations Based on Latent Purchase Motivations*. Rotterdam, Neth.: ERIM
- Jiang N, Li L. 2016. Doubly robust off-policy value evaluation for reinforcement learning. In *Proceedings of the 33rd International Conference on Machine Learning*, pp. 652–61. La Jolla, CA: Int. Mach. Learn. Soc.
- Kallus N. 2017. Balanced policy evaluation and learning. arXiv:1705.07384 [stat.ML]
- Keane MP. 2013. Panel data discrete choice models of consumer demand. In *The Oxford Handbook of Panel Data*, ed. BH Baltagi, pp. 54–102. Oxford, UK: Oxford Univ. Press

- Kitagawa T, Tetenov A. 2015. *Who should be treated? Empirical welfare maximization methods for treatment choice*. Tech. Rep., Cent. Microdata Methods Pract., Inst. Fiscal Stud., London
- Knox SW. 2018. *Machine Learning: A Concise Introduction*. Hoboken, NJ: Wiley
- Krizhevsky A, Sutskever I, Hinton GE. 2012. Imagenet classification with deep convolutional neural networks. In *Advances in Neural Information Processing Systems*, Vol. 25, ed. Z Ghahramani, M Welling, C Cortes, ND Lawrence, KQ Weinberger, pp. 1097–105. San Diego, CA: Neural Inf. Process. Syst. Found.
- Künzel S, Sekhon J, Bickel P, Yu B. 2017. Meta-learners for estimating heterogeneous treatment effects using machine learning. arXiv:1706.03461 [math.ST]
- Lai TL, Robbins H. 1985. Asymptotically efficient adaptive allocation rules. *Adv. Appl. Math.* 6:4–22
- LeCun Y, Bengio Y, Hinton G. 2015. Deep learning. *Nature* 521:436–44
- Levy O, Goldberg Y. 2014. Neural word embedding as implicit matrix factorization. In *Advances in Neural Information Processing Systems*, Vol. 27, ed. Z Ghahramani, M Welling, C Cortes, ND Lawrence, KQ Weinberger, pp. 2177–85. San Diego, CA: Neural Inf. Process. Syst. Found.
- Li L, Chen S, Kleban J, Gupta A. 2014. Counterfactual estimation and optimization of click metrics for search engines: a case study. In *Proceedings of the 24th International Conference on the World Wide Web*, pp. 929–34. New York: ACM
- Li L, Chu W, Langford J, Moon T, Wang X. 2012. An unbiased offline evaluation of contextual bandit algorithms with generalized linear models. In *Proceedings of 4th ACM International Conference on Web Search and Data Mining*, pp. 297–306. New York: ACM
- Matzkin RL. 1994. Restrictions of economic theory in nonparametric methods. In *Handbook of Econometrics*, Vol. 4, ed. R Engle, D McFadden, pp. 2523–58. Amsterdam: Elsevier
- Matzkin RL. 2007. Nonparametric identification. In *Handbook of Econometrics*, Vol. 6B, ed. J Heckman, E Learner, pp. 5307–68. Amsterdam: Elsevier
- Mazumder R, Hastie T, Tibshirani R. 2010. Spectral regularization algorithms for learning large incomplete matrices. *J. Mach. Learn. Res.* 11:2287–322
- Meinshausen N. 2007. Relaxed lasso. *Comput. Stat. Data Anal.* 52:374–93
- Mikolov T, Chen K, Corrado GS, Dean J. 2013a. Efficient estimation of word representations in vector space. arXiv:1301.3781 [cs.CL]
- Mikolov T, Sutskever I, Chen K, Corrado GS, Dean J. 2013b. Distributed representations of words and phrases and their compositionality. In *Advances in Neural Information Processing Systems*, Vol. 26, ed. Z Ghahramani, M Welling, C Cortes, ND Lawrence, KQ Weinberger, pp. 3111–19. San Diego, CA: Neural Inf. Process. Syst. Found.
- Mikolov T, Yih W, Zweig G. 2013c. Linguistic regularities in continuous space word representations. In *Proceedings of the 2013 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pp. 746–51. New York: Assoc. Comput. Linguist.
- Miller A. 2002. *Subset Selection in Regression*. New York: CRC Press
- Mnih A, Hinton GE. 2007. Three new graphical models for statistical language modelling. In *International Conference on Machine Learning*, pp. 641–48. La Jolla, CA: Int. Mach. Learn. Soc.
- Mnih A, Kavukcuoglu K. 2013. Learning word embeddings efficiently with noise-contrastive estimation. In *Advances in Neural Information Processing Systems*, Vol. 26, ed. Z Ghahramani, M Welling, C Cortes, ND Lawrence, KQ Weinberger, pp. 2265–73. San Diego, CA: Neural Inf. Process. Syst. Found.
- Mnih A, Teh YW. 2012. A fast and simple algorithm for training neural probabilistic language models. In *Proceedings of the 29th International Conference on Machine Learning*, pp. 419–26. La Jolla, CA: Int. Mach. Learn. Soc.
- Morris CN. 1983. Parametric empirical Bayes inference: theory and applications. *J. Am. Stat. Assoc.* 78:47–55
- Mullainathan S, Spiess J. 2017. Machine learning: an applied econometric approach. *J. Econ. Perspect.* 31:87–106
- Nie X, Wager S. 2019. Quasi-oracle estimation of heterogeneous treatment effects. arXiv:1712.04912 [stat.ML]
- Pennington J, Socher R, Manning CD. 2014. GloVe: global vectors for word representation. In *Proceedings of the 2014 Conference on Empirical Methods on Natural Language Processing*, pp. 1532–43. New York: Assoc. Comput. Linguist.



- Robins J, Rotnitzky A. 1995. Semiparametric efficiency in multivariate regression models with missing data. *J. Am. Stat. Assoc.* 90:122–29
- Rosenbaum PR, Rubin DB. 1983. The central role of the propensity score in observational studies for causal effects. *Biometrika* 70:41–55
- Ruiz FJ, Athey S, Blei DM. 2017. SHOPPER: a probabilistic model of consumer choice with substitutes and complements. arXiv:1711.03560 [stat.ML]
- Rumelhart DE, Hinton GE, Williams RJ. 1986. Learning representations by back-propagating errors. *Nature* 323:533–36
- Schapire RE, Freund Y. 2012. *Boosting: Foundations and Algorithms*. Cambridge, MA: MIT Press
- Scholkopf B, Smola AJ. 2001. *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*. Cambridge, MA: MIT Press
- Scott SL. 2010. A modern Bayesian look at the multi-armed bandit. *Appl. Stoch. Models Bus. Ind.* 26:639–58
- Semenova V, Goldman M, Chernozhukov V, Taddy M. 2018. Orthogonal ML for demand estimation: high dimensional causal inference in dynamic panels. arXiv:1712.09988 [stat.ML]
- Strehl A, Langford J, Li L, Kakade S. 2010. Learning from logged implicit exploration data. In *Advances in Neural Information Processing Systems*, Vol. 23, ed. Z Ghahramani, M Welling, C Cortes, ND Lawrence, KQ Weinberger, pp. 2217–25. San Diego, CA: Neural Inf. Process. Syst. Found.
- Sutton RS, Barto AG. 1998. *Reinforcement Learning: An Introduction*. Cambridge, MA: MIT Press
- Swaminathan A, Joachims T. 2015. Batch learning from logged bandit feedback through counterfactual risk minimization. *J. Mach. Learn. Res.* 16:1731–55
- Thomas P, Brunskill E. 2016. Data-efficient off-policy policy evaluation for reinforcement learning. In *Proceedings of the International Conference on Machine Learning*, pp. 2139–48. La Jolla, CA: Int. Mach. Learn. Soc.
- Thompson WR. 1933. On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika* 25:285–94
- Tibshirani R. 1996. Regression shrinkage and selection via the lasso. *J. R. Stat. Soc. B* 58:267–88
- Tibshirani R, Hastie T. 1987. Local likelihood estimation. *J. Am. Stat. Assoc.* 82:559–67
- van der Laan MJ, Rubin D. 2006. Targeted maximum likelihood learning. *Int. J. Biostat.* 2(1):34–56
- Van der Vaart AW. 2000. *Asymptotic Statistics*. Cambridge, UK: Cambridge Univ. Press
- Vapnik V. 2013. *The Nature of Statistical Learning Theory*. Berlin: Springer
- Varian HR. 2014. Big data: new tricks for econometrics. *J. Econ. Perspect.* 28:3–28
- Vilnis L, McCallum A. 2015. Word representations via Gaussian embedding. arXiv:1412.6623 [cs.CL]
- Wager S, Athey S. 2017. Estimation and inference of heterogeneous treatment effects using random forests. *J. Am. Stat. Assoc.* 113:1228–42
- Wan M, Wang D, Goldman M, Taddy M, Rao J, et al. 2017. Modeling consumer preferences and price sensitivities from large-scale grocery shopping transaction logs. In *Proceedings of the 26th International Conference on the World Wide Web*, pp. 1103–12. New York: ACM
- White H. 1992. *Artificial Neural Networks: Approximation and Learning Theory*. Oxford, UK: Blackwell
- Wooldridge JM. 2010. *Econometric Analysis of Cross Section and Panel Data*. Cambridge, MA: MIT Press
- Wu X, Kumar V, Quinlan JR, Ghosh J, Yang Q, et al. 2008. Top 10 algorithms in data mining. *Knowl. Inform. Syst.* 14:1–37
- Zeileis A, Hothorn T, Hornik K. 2008. Model-based recursive partitioning. *J. Comput. Graph. Stat.* 17:492–514
- Zhou Z, Athey S, Wager S. 2018. Offline multi-action policy learning: generalization and optimization. arXiv:1810.04778 [stat.ML]
- Zou H, Hastie T. 2005. Regularization and variable selection via the elastic net. *J. R. Stat. Soc. B* 67:301–20
- Zubizarreta JR. 2015. Stable weights that balance covariates for estimation with incomplete outcome data. *J. Am. Stat. Assoc.* 110:910–22